

Allgemeine Relativitätstheorie

Jörg Frauendiener*
Institut für Theoretische Astrophysik

WS/SS 2004/2005

Version: 4. Oktober 2005

*Since this script is available online I would appreciate any feedback from whoever reads it. Please send comments, suggestions, corrections to joerg.frauendiener@uni-tuebingen.de.

Inhaltsverzeichnis

1 Übersicht	7
2 Das Äquivalenzprinzip	11
3 Tensor-Kalkül	17
3.1 Multilineare Algebra und Tensoren	17
3.1.1 Vektoren und Kovektoren	17
3.1.2 Lineare Abbildungen	20
3.1.3 Tensoren	21
3.2 Euklidische Geometrie	27
4 Lorentz-Geometrie und Minkowski-Raum	33
4.1 Die relativistische Raumzeit	33
4.2 Physikalische Interpretation	36
5 Die gekrümmte Raumzeit	47
5.1 Die Raumzeit als Mannigfaltigkeit	47
5.1.1 Mannigfaltigkeiten	48
5.1.2 Der Tangentialraum	50
5.2 Tensorfelder	53
5.3 Der affine Zusammenhang	56
5.3.1 Motivation	57
5.3.2 Die kovariante Ableitung	59
5.3.3 Parallel-Transport	63
5.4 Krümmung	65
5.5 Riemannsche Geometrie	68
6 Die Feldgleichungen der Gravitation	71
6.1 Motivation	71
6.2 Der Energie-Impulstensor	77
6.3 Einstein-Tensor und Feldgleichung	79

7	Die Schwarzschild-Lösung	82
7.1	Herleitung der Schwarzschild-Metrik	82
7.2	Geodäten in der Schwarzschild-Metrik	87
7.3	Lichtstrahlen in der Schwarzschild-Metrik	91
7.4	Ein schwarzes Loch	95
8	Kugelsymmetrische statische Materieverteilungen	100
8.1	Die grundlegenden Gleichungen	100
8.2	Lösung der Gleichungen	103
8.3	Diskussion der Lösungen	108
9	Kosmologische Lösungen der Feldgleichungen	113
9.1	Das kosmologische Prinzip	114
9.2	Die Geometrie des Weltalls	116
9.3	Homogene und isotrope Mannigfaltigkeiten	119
9.4	Kinematische Eigenschaften von kosmologischen Modellen	124
9.5	Die Dynamik des Universums	135
9.5.1	Kosmologische Modelle mit $\Lambda = 0$	137
9.5.2	Kosmologische Modelle mit $\Lambda \neq 0$	142
9.5.3	Qualitative Diskussion der Friedmann-Lösungen	145
9.5.4	Die deSitter-Lösung	146
9.6	Eine kurze Geschichte des Universums	147
10	Eigenschaften von Schwarzen Löchern	152
10.1	Die Kruskal-Raumzeit	152
10.2	Die Einstein-Rosen Brücke	159
10.3	Die Kerr-Lösung	162
10.3.1	Null-Tetraden	162
10.3.2	Die Kerr-Metrik in Boyer-Lindquist Koordinaten	165
10.3.3	Kerr-Schild Koordinaten	167
10.3.4	Grundlegende Eigenschaften der Kerr-Metrik	168
10.3.5	Der Penrose-Prozess	172
A	Differentialgeometrische Grundlagen	175
A.1	Mannigfaltigkeiten und Tangentialraum	175
A.1.1	Mannigfaltigkeiten	176
A.1.2	Abbildungen zwischen Mannigfaltigkeiten	178
A.2	Felder	179
A.2.1	Skalarfelder	179
A.2.2	Kontravariante Vektorfelder	180
A.2.3	Kovariante Vektorfelder	185
A.2.4	Tensorfelder	188
A.3	Zusammenhang, kovariante Ableitung	190
A.4	Der Riemann-Tensor, Krümmung	198

Literaturverzeichnis

- [1] H. Bondi, *Relativity and Common Sense* (Dover, New York, 1980).
- [2] G. Börner, *The early universe. Texts and Monographs in Physics* (Springer-Verlag, Berlin, 1993).
- [3] R. d'Inverno, *Introducing Einstein's relativity* (Oxford University Press, Oxford, 1993).
- [4] A. Einstein (1905). *Zur Elektrodynamik bewegter Körper*. *Annalen der Physik* **17**, 891.
- [5] A. Einstein, *Relativity: The Special and the General Theory* (Crown Publisher, New York, 1961). Dazu gibt es auch ein deutsches Original.
- [6] R. Feynman, R. Leighton und M. Sands, *Lectures on Physics*, vol. 1 (Addison-Wesley, Reading, Massachusetts, 1977).
- [7] H. Gönner, *Einführung in die spezielle und allgemeine Relativitätstheorie* (Spektrum Akademie Verlag, Heidelberg, 1996).
- [8] S. Hawking und G. Ellis, *The large scale structure of space-times* (Cambridge University Press, Cambridge, 1973).
- [9] L. Hughston und P. Tod, *An Introduction to General Relativity* (Cambridge University Press, Cambridge, 1990).
- [10] H. Kaul (1996). *Skript zur ART Vorlesung 1996*. Tech. rep., Mathematisches Institut, Tübingen.
- [11] H. A. Lorentz, A. Einstein und H. Minkowski, *Das Relativitätsprinzip; eine Sammlung von Abhandlungen* (B. G. Teubner, Stuttgart, 1974). Dieses Buch enthält alle wesentlichen Originalartikel zur speziellen und allgemeinen Relativitätstheorie.
- [12] C. Misner, K. Thorne und J. Wheeler, *Gravitation* (Freeman, San Francisco, 1973).
- [13] W. Pauli, *Relativitätstheorie* (Springer-Verlag, Berlin, 2000).
- [14] R. Penrose und W. Rindler, *Spinors and Spacetime*, vol. 1 (Cambridge University Press, Cambridge, 1984).
- [15] R. Penrose und W. Rindler, *Spinors and Spacetime*, vol. 2 (Cambridge University Press, Cambridge, 1986).

- [16] W. Rindler, *Essential Relativity* (Springer-Verlag, Berlin, 1977), 2nd ed. Hier gibt es eine kürzlich bei Oxford University Press erschienene Neuauflage.
- [17] H. Ruder und M. Ruder, *Spezielle Relativitätstheorie* (Vieweg, Braunschweig, 1993).
- [18] R. Sexl und H. K. Schmidt, *Raum-Zeit-Relativität* (Vieweg, Braunschweig, 1978).
- [19] R. Sexl und H. Urbantke, *Gruppen, Teilchen, Relativität* (Springer-Verlag, Wien, 1992), 3rd ed.
- [20] N. Straumann, *General Relativity, with applications to astrophysics* (Springer-Verlag, 2004).
- [21] R. C. Tolman, *Relativity, Thermodynamics and Cosmology* (Dover, New York, 1987).
- [22] R. M. Wald, *General Relativity* (Chicago University Press, 1984).
- [23] H. Weyl, *Raum – Zeit – Materie* (Springer-Verlag, Berlin, 1988), 7th ed.
- [24] C. Will, *Theory and experiment in gravitational physics* (Cambridge University Press, 1981). In Neuauflage erschienen.
- [25] C. Will, *Was Einstein right?* (Oxford University Press, 1986).

1 Übersicht

Die *Allgemeine Relativitätstheorie* ist Einsteins Theorie der Gravitationskraft. Sie ist aber viel mehr als das. Im Gegensatz zu Newtons Gravitationstheorie, welche die Bewegung von Körpern unter dem Einfluss ihrer gegenseitigen Schwerkraft in einem festen, unveränderlichen, *absoluten Raum* und einer festen, *absoluten Zeit* beschreibt, handelt es sich bei der ART auch um eine Theorie von Raum und Zeit. Diese beiden sind nicht länger reine Kulissen, in denen die Phänomene der Physik stattfinden, sondern in der Einsteinschen Gravitationstheorie sind sie ein aktiver Mitspieler in diesem Stück.

Als Theorie von Raum und Zeit, bzw. wie sich herausstellen wird, als Theorie der Raum-Zeit, ist die ART natürlich eine Theorie vom Universum als Ganzen. Sie ist daher auch die Theorie, die für die Kosmologie, das Woher und Wohin unseres Universums zuständig ist. Dies ist sicherlich ein Grund für die Bedeutung dieser Theorie. Letztlich spielt sich jedes physikalische Phänomen in Raum und Zeit ab, ist also den Gesetzen der ART zu unterwerfen. Die ART ist also in diesem Sinne als *universell* zu betrachten.

Dies ist zwar aus physikalischer Sicht ein wesentlicher Grund, sich mit dieser Theorie auseinander zu setzen. Es ist aber sicherlich nicht der einzige Grund, warum sie eine Faszination auf viele Menschen ausübt. Die ART ist die erste Theorie, die allein aufgrund von theoretischen Überlegungen das Licht der Welt erblickt hat und erst im Nachhinein experimentell glänzend bestätigt wurde. Sie ist heute die am besten experimentell verifizierte physikalische Theorie 'auf dem Markt'. Die ART ist von einer berückenden Ästhetik, die sich dem Lernenden jedoch leider erst nach einem mühsamen Studium der mathematischen Grundlagen in vollem Ausmaß darstellt: Es passt alles!

Der wesentliche Schritt hin zur Speziellen Relativitätstheorie besteht in dem revolutionären Wandel unserer Vorstellung von Raum und Zeit, den sie von uns verlangt. Nach Einsteins Artikel über die Spezielle Relativitätstheorie [4] von 1905 war es Minkowski, der erkannte, dass Raum und Zeit sich aufgrund der Lorentz-Transformation zu einem vierdimensionalen 'Objekt' verschmelzen lassen. Diese *invariante Raum-Zeit* liefert eine vereinheitlichende Sichtweise auf Raum und Zeit. In seinen Arbeiten kommen solche Begriffe wie *Ereignis*, *Vierer-Vektor*, *Weltlinien*, *Lichtkegel* vor; alles Vorstellungen, die man in die ART übertragen kann, und die dort bis heute grundlegend für die Formulierung und Interpretation der Theorie sind.

Die ART geht jedoch wesentlich weiter. Sie schließt in die Beschreibung von Raum und Zeit auch *Gravitationsfelder* ein. Dies hat aber eine drastische Konsequenz: Gravitationsfelder sind veränderlich in Raum und Zeit. Dies bedeutet, dass eine konsistente

Beschreibung der Raum-Zeit sich nun auch von Ereignis zu Ereignis ändern kann. Die korrekte mathematische Beschreibung ist die *Differential-Geometrie*, bzw. eine spezielle Abart davon, die *Riemannsche Geometrie*.

Die differentialgeometrische Beschreibung beruht auf dem Postulat, die physikalische Raum-Zeit sei eine 4-dimensionale Mannigfaltigkeit, also ein abstraktes Objekt, dessen 'Punkte', die physikalischen Ereignisse, sich durch Angabe von vier reellen Zahlen in einem gewissen Rahmen ein-eindeutig charakterisieren lassen. Im allgemeinen wird man mehrere solcher *Koordinatensysteme* brauchen um ein solches Objekt in Gänze zu beschreiben. Die Vorstellung dabei ist die, dass dieses Objekt, die Mannigfaltigkeit, durch eine ganze Familie von solchen, sich teilweise überlappenden, Koordinatensystemen überdeckt wird.¹ Die grundlegenden Strukturen der Mannigfaltigkeit, ihre Topologie (ihre Gestalt) und ihre differentielle Struktur, d.h. die Vorschriften, wie man die infinitesimale Änderung von Funktionen beim Übergang zu infinitesimal benachbarten Punkten zu beschreiben hat, sind bestimmt durch die Relationen, die sich auf den Überlappgebieten zweier sich überlappenden Koordinatensystemen ergeben.

Viele physikalische Größen werden mathematisch durch *Tensoren* beschrieben. Dies sind Verallgemeinerungen von *Vektoren*. Entsprechend der Vektoralgebra gibt es eine Tensoralgebra, auch *Tensorkalkül* genannt, die es uns gestattet mit Tensoren zu rechnen. Dieser Kalkül lässt sich in natürlicher Weise auf Mannigfaltigkeiten übertragen. Hier tritt nun zunächst ein Problem auf. Dieser so definierte Tensorkalkül gilt jeweils nur *an einem Punkt*. Tensoren an verschiedenen Raum-Zeit-Punkten haben zunächst nichts miteinander zu tun. Will man Beziehungen zwischen zwei solchen Tensoren herstellen, z.B. die infinitesimale Änderung eines Tensors beim Übergang von einem Punkt zu einem infinitesimal benachbarten Punkt (also Tensoren differenzieren), dann muss man einen *Zusammenhang* zwischen den Raum-Zeit-Punkten herstellen. Dies ist eine wohl-definierte mathematische Struktur, die man durch ein Schema Γ^i_{jk} von 64 Funktionen charakterisieren kann. Mit Hilfe dieser *Christoffel-Symbole* lässt sich ein Begriff von Ableitung von Tensoren einführen, der unter anderem die Eigenschaft hat, dass Tensoren durch Ableiten wieder Tensoren ergeben. Man nennt diesen Prozess die *kovariante Ableitung*. Diese ist notwendig, wenn man physikalische Prozesse beschreiben will, die sich in Raum und Zeit vollziehen.

Die grundlegende Größe der ART ist die *Raum-Zeit-Metrik*. In der SRT werden die metrischen Eigenschaften von Raum und Zeit durch die *Minkowski-Metrik*

$$ds^2 = dt^2 - dx^2 - dy^2 - dz^2 = \sum_{i,k=0}^3 \eta_{ik} dx^i dx^k$$

beschrieben. Das wesentliche an diesem Ausdruck ist die Verteilung der Vorzeichen, die *Signatur* der Metrik ist $(1, -1, -1, -1)$, es handelt sich um eine Metrik mit *Lorentz-*

¹Man denke nur an die Oberfläche der Erdkugel, deren Punkte sich durch Angabe zweier Koordinaten (z. B. geografische Länge und Breite) eindeutig charakterisieren lassen. Man braucht mehrere *Karten*, gesammelt in einem *Atlas*, um die Erdoberfläche vollständig zu beschreiben.

Signatur. Diese Eigenschaft der Metrik ist es, die es uns erlaubt, von einem Raum-Zeit-Kontinuum zu sprechen, denn sie ist verantwortlich dafür, dass Lorentz-Transformationen in der SRT auftreten und damit solche Phänomene wie *Lorentz-Kontraktion* und *Zeitdilatation*.

In der ART ändert sich die Geometrie von Punkt zu Punkt und daher wird aus der konstanten Matrix

$$\eta_{ik} = \begin{pmatrix} 1 & & & \\ & -1 & & \\ & & -1 & \\ & & & -1 \end{pmatrix}$$

eine 4×4 -Matrix g_{ik} , von der nur verlangt wird, dass sie quadratisch und symmetrisch sei und Lorentz-Signatur besitze. Sie darf sich also von Punkt zu Punkt ändern. Aus der Minkowski-Metrik wird also eine Lorentz-Metrik

$$ds^2 = \sum_{i,k=0}^3 g_{ik} dx^i dx^k$$

Einsteins große Erkenntnis war das *Äquivalenzprinzip*, nämlich die Einsicht, dass sich Gravitationsfelder als die Angabe einer solchen Metrik auf dem Raum-Zeit-Kontinuum auffassen lassen und dass sich ihre Wirkung durch die *Krümmung* dieser Metrik beschreiben lässt. Die Tatsache, dass wir ein Raum-Zeit-Kontinuum beschreiben wollen, steckt allein in der Forderung der Lorentz-Signatur.

Wir haben es nun also mit zwei Größen zu tun, die gegeben sein müssen, wenn wir physikalische Prozesse in einer sich ändernden Raum-Zeit beschreiben wollen: Zusammenhang und Metrik, die von vornherein nichts miteinander zu tun haben. Ein Indiz für die Konsistenz der Theorie folgt nun aus der rein mathematischen Tatsache, dass es zu einer gegebenen Metrik genau einen ausgezeichneten Zusammenhang, eine eindeutige kovariante Ableitung, gibt. Das bedeutet, dass die Metrik allein die Verhältnisse zwischen verschiedenen Ereignissen bestimmt. Der Zusammenhang ist durch die Eigenschaft charakterisiert, dass die Metrik von Punkt zu Punkt im Sinne der kovarianten Ableitung konstant bleibt.

Nun haben wir also eine Metrik auf dem Raum-Zeit-Kontinuum. Wie aber ist die Metrik bestimmt? Gibt es ein Naturgesetz, welches die Metrik festlegt, und welche Größen beeinflussen diese Festlegung? Der Prozess zur Aufstellung dieses Gesetzes war extrem langwierig. Einstein brauchte mehrere Versuche dazu. Das Resultat aber kann sich sehen lassen.

Welche Größen können die Metrik, also die Struktur von Raum und Zeit, beeinflussen? Es bleibt nur die Materie, die sich in der Raum-Zeit befindet. In der ART beschreibt diese Materie durch einen sogenannten *Energie-Impuls-Tensor* T^{ab} , eine Größe, die durch die Art der Materie bestimmt ist. Die charakteristische Eigenschaft des Energie-Impuls-Tensors ist seine *Divergenzfreiheit*,

$$\nabla_a T^{ab} = 0.$$

Diese Eigenschaft drückt aus, dass Energie und Impuls in gewisser Weise erhalten sind.

Es ist nun eine weitere mathematische Tatsache, dass es genau einen Tensor G_{ab} gibt, den man aus der Metrik konstruieren kann, der also durch die Raum-Zeit-Struktur bestimmt ist und der ebenfalls notwendigerweise divergenzfrei ist. Dieser *Einstein-Tensor* ist ein Maß für die *Krümmung* der Raum-Zeit. Nun haben wir zwei Tensoren, von denen einer durch die Raum-Zeit-Struktur, der andere durch die Materie bestimmt ist. Beide sind divergenzfrei. Was liegt näher, als zu postulieren, dass beide proportional zueinander sind? Genau das hat Einstein getan und die Gleichung

$$G_{ab} = \frac{8\pi G}{c^4} T_{ab}$$

postuliert. Dies ist die *Einsteinsche Feldgleichung*, die die Metrik der Raum-Zeit bei gegebener Materieverteilung bestimmt. Diese Gleichung ist eine gewisse Verallgemeinerung des Newtonschen Gravitationsgesetzes in dem Sinne, dass sich dieses aus jener in einem geeigneten Grenzübergang ergibt. Sie ist aber viel mehr, weil sich die ganze Betrachtungsweise gegenüber der klassischen Physik geändert hat.

Gravitation
Fernwirkungsprinzip

Nachdem Einstein die Prinzipien der Mechanik so geändert hatte, dass sie mit der Maxwellschen Theorie des Elektromagnetismus kompatibel wurde, wäre es ein natürlicher Schritt gewesen, nun die zweite klassische Wechselwirkung, die Schwerkraft oder **Gravitation** mit den Prinzipien der Speziellen Relativitätstheorie zu vereinen. Die bisher akzeptierte Gravitationstheorie Newtons war mit der neuen Theorie nicht kompatibel, denn sie fußt auf einem **Fernwirkungsprinzip**: jede Störung 'hier' macht sich unmittelbar, sofort, an jedem anderen Ereignis 'dort' bemerkbar. Dies ist natürlich mit der Existenz einer maximalen Ausbreitungsgeschwindigkeit nicht vereinbar. Man könnte sich vorstellen, die Gravitationstheorie, die ja einem ähnlichen Gesetz wie dem Coulomb-Gesetz der Elektrostatik genügt, durch Hinzufügen entsprechender Terme in eine der Maxwelltheorie ähnliche Theorie umformen zu können.

Abgesehen davon, dass dieser Weg nicht zum Ziel führen würde, es war auch nicht der Weg, den Einstein einschlug. Zum Ausgangspunkt seiner Überlegungen nahm er die weithin bekannte, jedoch unverstandene Tatsache, dass 'alle Körper gleich schnell fallen'.

2 Das Äquivalenzprinzip

Von Galilei wird behauptet, er habe seine Fall-Experimente am schiefen Turm von Pisa durchgeführt. Ob dies tatsächlich der Fall war, sei dahin gestellt. Wesentlich für uns ist die Einsicht, zu der er durch seine Experimente gekommen ist:

In einem Gravitationsfeld fallen alle Körper gleich schnell.

Eine andere, etwas präzisere, Formulierung der gleichen Tatsache ist: Zwei Teilchen, die am gleichen Ort zur gleichen Zeit mit der gleichen Anfangsgeschwindigkeit ausgestattet werden, durchlaufen die gleiche Bewegung unabhängig von ihrer inneren Zusammensetzung. Die Gravitationskraft beeinflusst alle Teilchen gleich.

Galileis Einsicht hat in der klassischen Mechanik von Newton keine tiefere Bedeutung. Sie wird sozusagen als zusätzliches 'Schmankerl' mitgenommen. Es dauerte ungefähr 300 Jahre, bis ihre tiefe Bedeutung gewürdigt wurde. Sie ist nun ein wesentlicher Teil des Fundaments, auf dem die Einsteinsche Gravitationstheorie – und damit die gesamte moderne Physik – aufgebaut ist, denn es ist auf Grund dieser Einsicht, dass wir unsere Raumzeit als eine im allgemeinen gekrümmte Mannigfaltigkeit zu beschreiben haben. Man nennt diese Einsicht von Galilei auch das **schwache Äquivalenzprinzip**.

schwache
Äquivalenzprinzip

Nach Newton wird die Bewegung eines Körpers K_1 unter der Wirkung der Schwerkraft durch zwei Gesetze bestimmt. Da ist zunächst das zweite Newtonsche Gesetz der Mechanik, dem gemäß die Beschleunigung des Körpers proportional ist zur einwirkenden Kraft

$$\mathbf{F} = m_t \mathbf{a}.$$

Der Proportionalitätsfaktor ist ein Maß für den Widerstand, den der Körper einer Beschleunigung entgegensetzt, also für die Trägheit des Körpers. Er wird **träge Masse** genannt. Die dabei wirkende Kraft wird durch das Newtonsche Gravitationsgesetz bestimmt

träge Masse

$$\mathbf{F} = -\frac{G m_s M}{r^3} \mathbf{r},$$

wo M die Masse eines anderen Körpers K_2 ist. Der Vektor $\mathbf{r}$ ist der Verbindungsvektor zwischen den beiden Körpern $\mathbf{r} = \mathbf{r}_2 - \mathbf{r}_1$. Die Masse m_s nennt man die **passive schwere Masse** des Körpers, weil sie angibt, wie der Körper auf das von M erzeugte Gravitationsfeld reagiert.

passive schwere
Masse

Übung 2.1: Man nennt M auch oft die *aktive schwere Masse*, weil sie die Kraft auf den Körper K_1 verursacht. Warum sind aktive und passive Masse eines Körpers gleich?

Es gibt keinen Grund innerhalb der Newtonschen Theorie anzunehmen, dass diese beiden Massenbegriffe gleich sind. Es wäre durchaus vorstellbar, dass die schwere Masse, d.h. die Fähigkeit eines Körpers eine Kraft auf andere Körper zu erzeugen von seinem Material, seiner inneren Zusammensetzung abhängen könnte, und dies in einer anderen Art und Weise wie die träge Masse. Für Newton stellte dies kein Problem dar. Er nahm die Gleichheit von träger und schwerer Masse als gottgegeben hin. Damit wird die Bewegungsgleichung des Körpers von $m_t = m_s$ unabhängig und damit erfährt jeder Körper unter gleichen Bedingungen die gleiche Beschleunigung, die Körper fallen gleich schnell. Dass dies nicht selbstverständlich ist, wird durch die Betrachtung von geladenen Teilchen deutlich. Diese unterliegen dem Coulomb-Gesetz, welches genau die gleiche Form wie das Gravitationsgesetz hat. Dennoch ist es hier nicht so, dass Teilchen unterschiedlicher Art, also möglicherweise unterschiedlicher Ladung die gleiche Bahn durchliefen. Vielmehr wird genau dieser Unterschied z.B. in Massenspektrographen ausgenutzt.

Im Gegensatz zur Newtonschen Theorie ist die Einsteinsche Gravitationstheorie auf dem 'Prinzip der Gleichheit von schwerer und träger Masse' aufgebaut. Es ist dieses Äquivalenzprinzip, welches zur Konsequenz hat, dass das Raumzeit-Kontinuum, in dem wir leben, gekrümmt ist. Die ART steht und fällt mit seiner Gültigkeit. Daher ist es beruhigend zu wissen, dass die experimentelle Evidenz für die Gültigkeit des schwachen Äquivalenzprinzips überwältigend ist. Durch verschiedenartige Experimente¹ (Fallexperimente, Torsionswaage, usw.) konnte das schwache Äquivalenzprinzip bis auf eine Genauigkeit von 10^{-13} bestätigt werden.

Wie kommt es, dass eine solch 'unschuldige' Aussage über das Fallverhalten von Körpern eine so tiefgreifende Konsequenz wie die Krümmung unserer Raumzeit zur Folge haben kann? In einem frei fallenden Fahrstuhl bewegen sich die Insassen schwerelos umher. Astronauten werden probenhalber der Schwerelosigkeit ausgesetzt, indem man sie in ein Flugzeug steckt, das ballistische Bahnen, also Bahnen eines frei fallenden Körpers, zu fliegt. Stellen wir uns schließlich einen Astronauten auf einem 'Raumspaziergang' ausserhalb des Spaceshuttles 'Atlantis' vor, die sich in einer Erdumlaufbahn befindet. Trotz des riesigen Planets ganz in der Nähe stürzen weder die Atlantis noch der Astronaut auf die Erde zu, sondern bewegen sich unbeeinflusst von Gravitationskräften schwerelos umher.

Diese drei Beispiele zeigen uns, dass wir, zumindest in einem kleinen Raumzeitgebiet, die Effekte der Gravitation 'wegtransformieren' können, indem wir ein geeignetes Bezugssystem wählen. Bezüglich des Fahrstuhls bewegen sich die Insassen kräftefrei, ebenso wie die Astronauten bezüglich des Flugzeugs und ebenso wie der Astronaut bezüglich der 'Atlantis'. Dabei führt aber jedes dieser Bezugssysteme selbst eine beschleunigte Bewegung bezüglich der Erde aus.

¹Eine ausführliche Beschreibung der verschiedenen Experimente, die zum Test des Äquivalenzprinzips wie auch anderer Aspekte der Theorie findet man in dem Buch: C. Will, *Theory and experiment in gravitational physics*, Cambridge University Press.

Der Effekt von Beschleunigung und Gravitation auf einen Körper sind wegen der Gleichheit von schwerer und träger Masse nicht unterscheidbar und deshalb 'gegeneinander aufrechenbar'. Man kann Beschleunigungs- durch Gravitationseffekte erzeugen und umgekehrt. Betrachten wir noch einmal das fallende Labor. Es führt bezüglich der Erde, welche wir für diese Betrachtung als ein Inertialsystem annehmen, eine beschleunigte Bewegung durch.

Übung 2.2: Zeigen Sie, dass in einem frei fallenden Bezugssystem die Gesetze der Newtonschen Mechanik unverändert gelten.

In diesem Sinne verhält sich das frei fallende System ebenfalls wie ein Inertialsystem. Die Gravitation wurde wegtransformiert. Einstein erhebt nun diese Tatsache zum **Äquivalenzprinzip**

Äquivalenzprinzip

Alle lokalen, frei fallenden und nicht rotierenden Laboratorien sind zur Beschreibung physikalischer Vorgänge vollständig äquivalent.

Das heißt, in all diesen Laboratorien verlaufen physikalische Phänomene genau gleich. Stellen wir uns ein solches 'Kabinenlabor' vor, welches sich weit weg von irgendwelchen Massen, kräftefrei bewegt. In diesem System gilt das erste Newtonsche Gesetz: Ein kräftefreier Körper bewegt sich auf einer geraden Linie. Das Kabinenlabor ist also ein Inertialsystem. Dem Äquivalenzprinzip zufolge ist dann aber jedes lokale, frei fallende und nicht rotierende System ein Inertialsystem. Es handelt sich also hierbei um ein lokales Inertialsystem.

Der Begriff des Inertialsystems wird durch das Äquivalenzprinzip zugleich erweitert und eingeschränkt. *Erweitert* in dem Sinne, dass Inertialsysteme relativ zueinander beschleunigt sein dürfen. Und *eingeschränkt* in dem Sinne, dass wir nur noch *lokale* Inertialsysteme betrachten dürfen. Denn es ist klar, dass man durch Beschleunigung eines ausgedehnten Labors die gleiche Wirkung erzeugt, wie ein *homogenes* Schwerfeld. Da es aber durchaus inhomogene Schwerfelder gibt (Schwerfeld der Erde), muss die Ausdehnung der Systeme beschränkt werden. Im Grenzfall reduziert sich ein solches frei fallendes System auf eine Weltlinie mit drei, an jedem Ereignis angehefteten Achsen, die auf infinitesimal benachbarte Ereignisse zeigen.

In Rahmen der SRT laufen kräftefreie Körper auf geraden Weltlinien. Wir hatten gesehen, dass diese Linien sich lokal dadurch charakterisieren lassen, dass die Änderung ihres Tangentenvektors an jedem Ereignis proportional zum Tangentenvektor selber ist. Sie sind also **autoparallele Kurven**. Das heißt, sie weichen nicht von ihrer Richtung ab, sie sind so gerade wie nur irgend möglich. Wenn wir nun ein Teilchen betrachten, das sich unter dem Einfluß eines Schwerfeldes bewegt, also frei fällt, dann ist es bezüglich seines (mitfallenden) Bezugssystems in Ruhe, also kräftefrei. Auf Grund des Äquivalenzprinzips muss auch jetzt die Weltlinie des Teilchens eine autoparallele Kurve sein. Damit haben wir das Bewegungsgesetz eines kräftefreien Massenpunkts gefunden

autoparallele
Kurven

Ein frei fallendes (Test-)Teilchen bewegt sich auf einer autoparallelen Kurve.

Die Beschränkung auf lokale Inertialsysteme ist tatsächlich notwendig. Erlauben wir uns nämlich, Messungen durchzuführen, die nicht ganz lokal sind, so lassen sich gravitative und Beschleunigungseffekte unterscheiden. Dazu betrachten wir eine Wolke von kleinen Teilchen im freien Raum, die auf die Erde zufallen (vgl. Abb. 2.1). Setzen wir

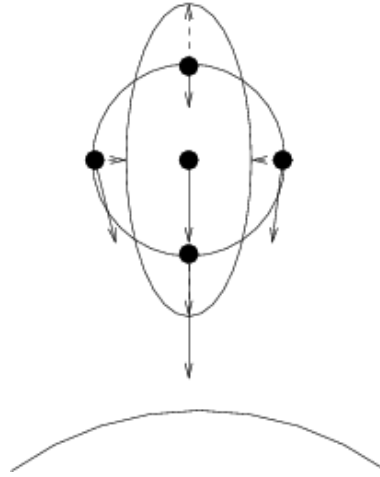


Abbildung 2.1: Zum Gezeiten-Effekt

uns auf das Ruhssystem eines im Zentrum der Wolke befindlichen Teilchens, dann fallen wir mit der Wolke auf die Erde zu. Da es sich um eine ausgedehnte Wolke handelt, fallen weiter entfernte Teilchen unterschiedlich schnell auf die Erde zu. Die Teilchen 'unter uns' fallen schneller, während die über uns langsamer fallen. Die Teilchen neben uns fallen zwar gleich schnell, aber auf Bahnen, die nicht parallel zur unseren sind, weil sie auch direkt auf den Erdmittelpunkt zeigen. Der Netto-Effekt dieser Bewegung ist also der, dass sich die Teilchen in vertikaler Richtung von uns entfernen, während die Teilchen neben uns näher kommen. Eine vormals kugelförmige Verteilung wird im Laufe der Zeit zu einem Ellipsoid verformt. Dies ist der sogenannte **Gezeiten-Effekt**.

Gezeiten-Effekt

Dieser gravitative Effekt lässt sich nicht wegtransformieren, da es sich hier um einen Effekt eines inhomogenen Schwerfeldes handelt. Es ist aber auch gleichzeitig der typische Effekt, den ein Schwerfeld erzeugen kann.

Betrachten wir ein idealisiertes Raumzeit-Diagramm dieser Situation (vgl. Abb. 2.2) Die Deformation eines Kreises in ein Ellipsoid erfolgt in der Raumzeit durch den freien Fall von Teilchen. Die Weltlinien der Teilchen sind also autoparallelen Kurven. Trotzdem ist es offensichtlich, dass es sich hierbei nicht um Geraden im üblichen Sinne handeln kann. Vielmehr sind die Kurven 'krumm'. Es muss sich um autoparallele Kurven in einer gekrümmten Raumzeit handeln.

Ein ähnliches Ergebnis liefert die folgende Betrachtung. In einem frei fallenden Labor (Fahrstuhl) bewege sich ein Lichtblitz senkrecht zur Fallrichtung. Gemäß dem Äquivalenzprinzip handelt es sich bei diesem System um ein lokales Inertialsystem, das Licht

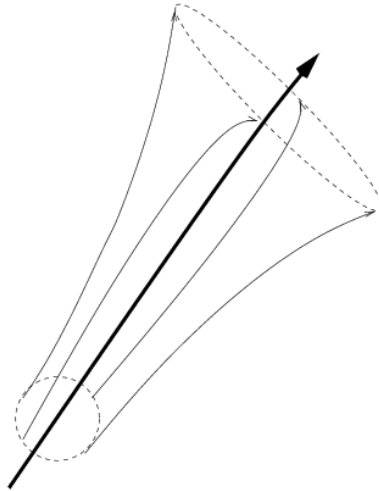


Abbildung 2.2: Raumzeit-Diagramm des Gezeiten-Effekts

breitet sich demnach gradlinig aus. Das bedeutet aber, dass sich dieses Licht relativ zur Erde auf einer Parabel bewegen muss. Wir kommen also zu dem Schluß, dass Licht in einem Gravitationsfeld abgelenkt wird. Dieses Argument beruht allein darauf, dass sich Licht mit endlicher Geschwindigkeit ausbreitet, also Zeit braucht um das Labor zu durchqueren, welches in der gleichen Zeit ein Stück weit fällt. D.h. allein die Tatsache, dass die Lichtgeschwindigkeit endlich ist, bedeutet, dass Licht ein 'Gewicht' hat, in dem Sinne, dass es auf die Erde zufällt.

Eine weitere Konsequenz des Äquivalenzprinzips ist die **gravitative Rot- bzw. Blauverschiebung** des Lichts im Gravitationsfeld. Wir betrachten wieder ein frei fallendes Labor und nehmen an, dass sich in dem Moment, in dem sich das Labor in Bewegung setzt, das Licht an der Decke des Labors angeschaltet wird und sich nach unten ausbreitet, wo es von einem Beobachter A registriert wird. Da das Labor frei fällt, handelt es sich um ein Inertialsystem und der Beobachter A stellt keine Frequenzverschiebung fest. Nehmen wir nun an, dass das Labor an einem Beobachter B vorbeifällt, welcher bezüglich der Erde in Ruhe ist. Von A aus betrachtet, bewegt sich B dem Licht entgegen. B muss also aufgrund des Dopplereffekts eine Frequenzverschiebung des Lichts feststellen. Da sich die Frequenz für A nicht ändert, sieht B eine Blauverschiebung. Er schließt also aus dieser Beobachtung, *dass Licht im Gravitationsfeld frequenzverschoben wird*, und zwar in Richtung Blau, wenn es in das Gravitationsfeld hinein läuft, bzw. Richtung Rot, wenn es heraus kommt.

gravitative Rot- bzw.
Blauverschiebung

Das bedeutet aber auch, dass die Periode des Lichts im Bereich starker Gravitationsfelder kleiner ist, als im Bereich schwacher Felder. Da das Licht ein periodischer Vorgang ist, können wir es als eine Uhr betrachten, die mit der Lichtfrequenz 'tickt'. Um dies etwas genauer zu sehen, betrachten wir noch einmal zwei Beobachter A und B, die sich relativ zueinander in Ruhe im Gravitationsfeld der Erde befinden. A habe eine Atom-

uhr bei sich, die mit einer Frequenz ν tickt, d.h. sie erzeugt Licht mit dieser Frequenz. Der Beobachter A befinde sich 'unterhalb' von B, also im stärkeren Gravitationsfeld. Nehmen wir an, die Uhr bei A schlage N mal, d.h. sie möge einen Wellenzug mit N Maxima aussenden. Offensichtlich dauert es eine Zeit $T_A = N/\nu$, bis dieser Wellenzug ausgesandt ist. Nun bewegt sich das Licht aus dem Gravitationsfeld heraus und wird daher rotverschoben. Es kommt also bei B mit einer Frequenz $\nu_B = \nu - \Delta\nu < \nu$ bei B an und B benötigt eine Zeit $T_B = N/\nu_B$, bis er alle N Maxima registriert hat. Nun ist

$$N = \nu T_A = \nu_B T_B = \nu \left(1 - \frac{\Delta\nu}{\nu}\right) T_B.$$

Also ist $T_A = \left(1 - \frac{\Delta\nu}{\nu}\right) T_B < T_B$, d.h. auf der Uhr von A vergeht weniger Zeit als bei B. Die Uhr bei A 'geht nach', die Zeit in A vergeht langsamer als bei B.

Diese qualitativen Überlegungen zeigen uns, dass wir die physikalische Raumzeit nicht mehr als den einfachen affinen Raum der SRT modellieren dürfen. Wir sehen, dass der Begriff eines Inertialsystems nur noch lokalen Sinn hat. Eine gerade Linie ist nicht notwendigerweise eine Gerade und Licht breitet sich nicht auf Geraden aus.

Wir brauchen also eine Beschreibung der Raumzeit, die es gestattet der Äquivalenz der verschiedenen lokalen Inertialsysteme Rechnung zu tragen und die keine 'a priori' Annahmen über die globale Struktur der Raumzeit macht. Eine solche Theorie ist die Theorie der Mannigfaltigkeiten und insbesondere die Differentialgeometrie, der wir uns in dieser Vorlesung zuwenden müssen. Zunächst aber benötigen wir einen Formalismus, der es gestattet, physikalische Gesetzmäßigkeiten unabhängig vom Bezugssystem zu formulieren. Dies führt uns auf die Theorie der Tensoren.

3 Tensor-Kalkül

3.1 Multilineare Algebra und Tensoren

Eine wesentliche Voraussetzung bei der Formulierung der ART ist die **Tensorrechnung**. Alle physikalischen Größen haben einen tensoriellen Charakter und die physikalischen Gesetzmäßigkeiten lassen sich infolgedessen als **Tensorgleichungen** formulieren. Die wesentliche Eigenschaft dieser Objekte ist ihre **Invarianz** unter Koordinatentransformationen. Was dies genau heißt, wird im Laufe der Vorlesung erläutert werden. Hier werden zunächst die algebraischen Eigenschaften dieser Objekte diskutiert. Ziel dieses Abschnitts ist nicht, eine mathematisch einwandfreie und lückenlose Darstellung zu geben, sondern vielmehr eine gewisse Gewöhnung und Rechenfertigkeit mit diesen Objekten zu erlangen.

Tensorrechnung
Tensorgleichungen
Invarianz

3.1.1 Vektoren und Kovektoren

Wir betrachten einen endlich dimensionalen Vektorraum $\mathbb{V}$ über einem Körper, der beliebig sein kann, den wir aber der Einfachheit halber als den Körper $\mathbb{R}$ der reellen Zahlen annehmen. Vektoren kann man addieren und mit Skalaren (reellen Zahlen) multiplizieren, also Linearkombinationen

$$r\mathbf{v} + s\mathbf{u}$$

mit $r, s \in \mathbb{R}$ und $\mathbf{u}, \mathbf{v} \in \mathbb{V}$ bilden. Ist $\dim_{\mathbb{R}} \mathbb{V} = n$, dann gibt es ein n -tupel $(\mathbf{e}_1, \mathbf{e}_2, \dots, \mathbf{e}_n)$ von linear unabhängigen Vektoren, eine **Basis**, so dass sich jedes $\mathbf{v} \in \mathbb{V}$ als *eindeutige* Linearkombination in diesen Vektoren schreiben lässt

Basis

$$\mathbf{v} = v^1 \mathbf{e}_1 + v^2 \mathbf{e}_2 + \dots + v^n \mathbf{e}_n = \sum_{i=1}^n v^i \mathbf{e}_i.$$

Die Skalare v^i nennt man die **Koordinaten** des Vektors $\mathbf{v}$ bzgl. der Basis $\mathbf{e}_i$. Basen in einem Vektorraum sind nicht eindeutig. Ist $(\hat{\mathbf{e}}_i)_{i=1:n}$ eine weitere Basis, dann muss offensichtlich gelten

Koordinaten

$$\hat{\mathbf{e}}_k = \sum_{i=1}^n s^i_k \mathbf{e}_i \quad (3.1.1)$$

mit n^2 Skalaren s^i_k . Da hier eine Basis auf eine andere Basis abgebildet wird, müssen die Skalare eine reguläre $n \times n$ -Matrix bilden, d.h. $\det(s^i_k) \neq 0$. Und für einen beliebigen Vektor $\mathbf{v}$ gilt beim **Basiswechsel**

Basiswechsel

$$\mathbf{v} = \sum_{k=1}^n \hat{v}^k \hat{\mathbf{e}}_k = \sum_{k=1}^n \hat{v}^k \left(\sum_{i=1}^n s^i_k \mathbf{e}_i \right) = \sum_{i=1}^n \left(\sum_{k=1}^n s^i_k \hat{v}^k \right) \mathbf{e}_i.$$

Daher gilt also die **Koordinatentransformation**

$$v^i = \sum_{k=1}^n s^i_k \hat{v}^k \quad i = 1, \dots, n \quad (3.1.2)$$

kontravariant

zwischen den Koordinaten des Vektors $\mathbf{v}$ beim Wechsel der Basis. Vergleichen wir (3.1.1) und (3.1.2) stellen wir fest, dass sich die Basisvektoren und die auf sie bezogenen Koordinaten 'gegenläufig' transformieren. Man nennt das Verhalten der Koordinatentransformation daher auch **kontravariant**. Also: die Koordinaten eines Vektors in Bezug auf eine Basis verhalten sich beim Basiswechsel kontravariant zur Transformation der Basisvektoren. Es ist wichtig, sich klar zu machen, dass es einen wesentlichen Unterschied gibt zwischen dem Vektor $\mathbf{v}$ und den Koordinaten v^i . Letztere haben nur einen Sinn *in Bezug auf eine Basis*, die benötigt wird, um den Vektor $\mathbf{v}$ zu rekonstruieren. Dieser ist unabhängig von der Wahl der Basis definiert. Wenn man dann an der Basis 'dreht', dann müssen sich entsprechend die Koordinaten 'drehen', so dass der Vektor selbst ungeändert bleibt. Genau dies ist der Effekt des kontravarianten Transformationsverhaltens.

Stellen wir die Linearkombination $r\mathbf{u} + s\mathbf{v}$ bezüglich einer Basis dar, ergibt sich

$$r\mathbf{u} + s\mathbf{v} = r \sum_{i=1}^n u^i \mathbf{e}_i + s \sum_{i=1}^n v^i \mathbf{e}_i = \sum_{i=1}^n (ru^i + sv^i) \mathbf{e}_i.$$

Bezüglich einer anderen Basis ergibt sich natürlich

$$r\mathbf{u} + s\mathbf{v} = \sum_{i=1}^n (r\hat{u}^i + s\hat{v}^i) \hat{\mathbf{e}}_i.$$

Dies zeigt uns zweierlei: erstens, die Linearkombination zweier Vektoren berechnet man, indem die Koordinaten der Vektoren bzgl. einer Basis entsprechend linear kombiniert, und zweitens, *diese Rechenvorschrift ist basisunabhängig*. Dieses ist eine grundlegende Einsicht, die im weiteren noch ausgebaut wird.

Es gibt zwei Standpunkte, um dies zu betrachten. Vom mathematischen Standpunkt aus ist dies eine Trivialität, denn der Vektor $r\mathbf{u} + s\mathbf{v}$ ist unabhängig von jeder Basis definiert, daher darf die Berechnung der Linearkombination nicht von der Wahl der Basis abhängen. Dass dies nicht der Fall ist, liegt genau an dem kontravarianten Transformationsverhalten von Basisvektoren und Vektor-Koordinaten.

Vom konstruktiven, physikalischen Standpunkt aus betrachtet, kann man sagen: sind zwei n -tupel u^i und v^i gegeben, und berechnet man ein neues n -tupel durch die Linearkombination $ru^i + sv^i$, dann erhält diese Rechenvorschrift eine basisunabhängige Bedeutung, *sofern man verlangt, dass die n -tupel sich beim Basiswechsel kontravariant transformieren*. Beide Auffassungen sind legitim und werden bei der weiteren Entwicklung der Theorie gebraucht.

Zu jedem Vektorraum $\mathbb{V}$ lässt sich in kanonischer Weise ein weiterer Vektorraum konstruieren. Eine Linearform α ist eine lineare Abbildung

$$\alpha : \mathbb{V} \rightarrow \mathbb{R}, \quad \mathbf{v} \mapsto \alpha(\mathbf{v}) =: \langle \alpha, \mathbf{v} \rangle,$$

die also jedem Vektor $\mathbf{v}$ eine Zahl $\langle \alpha, \mathbf{v} \rangle \in \mathbb{R}$ zuordnet. Mit der üblichen Addition und Skalarmultiplikation für Abbildungen nach $\mathbb{R}$ ausgestattet¹, bildet die Menge aller Linearformen auf $\mathbb{V}$ einen reellen Vektorraum, den (algebraischen) **Dualraum** $\mathbb{V}^*$. Die Elemente des Dualraums nennt man auch **Kovektoren**.

Dualraum
Kovektoren

Welche Dimension hat dieser Raum? Dazu nehmen wir eine Basis $(\mathbf{e}_i)_{i=1:n}$ in $\mathbb{V}$, beachten $\mathbf{v} = \sum_{i=1}^n v^i \mathbf{e}_i$ und konstruieren folgende Abbildungen

$$\begin{aligned} \omega^1 : \mathbf{v} &\mapsto \langle \omega^1, \mathbf{v} \rangle =: v^1 \\ \omega^2 : \mathbf{v} &\mapsto \langle \omega^2, \mathbf{v} \rangle =: v^2 \\ &\vdots \\ \omega^n : \mathbf{v} &\mapsto \langle \omega^n, \mathbf{v} \rangle =: v^n. \end{aligned}$$

Diese sind offensichtlich linear, liegen also in $\mathbb{V}^*$. Insbesondere gilt die **Dualitätsrelation** zwischen den beiden Basen

Dualitätsrelation

$$\langle \omega^i, \mathbf{e}_k \rangle = \delta_k^i = \begin{cases} 1 & i = k \\ 0 & i \neq k \end{cases}.$$

Ist α eine beliebige Linearform, dann können wir schreiben

$$\langle \alpha, \mathbf{v} \rangle = \left\langle \alpha, \sum_{i=1}^n v^i \mathbf{e}_i \right\rangle = \sum_{i=1}^n v^i \langle \alpha, \mathbf{e}_i \rangle = \sum_{i=1}^n \alpha_i \langle \omega^i, \mathbf{v} \rangle,$$

wobei wir mit $\alpha_i = \langle \alpha, \mathbf{e}_i \rangle$ das Bild der Basisvektoren unter α bezeichnet haben. Wir haben also das Resultat

$$\alpha = \sum_{i=1}^n \alpha_i \omega^i$$

¹ $\langle r\alpha + s\beta, \mathbf{v} \rangle := r \langle \alpha, \mathbf{v} \rangle + s \langle \beta, \mathbf{v} \rangle$ für alle $r, s \in \mathbb{R}$, $\alpha, \beta \in \mathbb{V}^*$ und $\mathbf{v} \in \mathbb{V}$

d.h., die ω^i erzeugen $\mathbb{V}^*$. Andererseits sind sie aber auch linear unabhängig, wie man leicht an der folgenden Überlegung sieht:

$$\begin{aligned} 0 = \sum_{i=1}^n r_i \omega^i &\iff 0 = \left\langle \sum_{i=1}^n r_i \omega^i, \mathbf{e}_k \right\rangle \quad \text{für alle } k = 1 : n \\ &\iff 0 = \sum_{i=1}^n r_i \delta_k^i = r_k \quad \text{für alle } k = 1 : n. \end{aligned}$$

duale Basis Damit bilden diese Linearformen eine Basis von $\mathbb{V}^*$, die zu $\mathbf{e}_k$ **duale Basis**².

Wählen wir eine andere Basis $\hat{\mathbf{e}}_k$ in $\mathbb{V}$, dann gibt es auch eine andere, dazu duale Basis $\hat{\omega}^i$. Es ist nun leicht, die folgende Beziehung zwischen den dualen Basen beim Basiswechsel in $\mathbb{V}$ herzuleiten

$$\hat{\mathbf{e}}_k = \sum_{i=1}^n s_k^i \mathbf{e}_i \iff \omega^i = \sum_{k=1}^n s_k^i \hat{\omega}^k.$$

Jeder Kovektor $\alpha \in \mathbb{V}^*$ hat die Darstellung

$$\alpha = \sum_{i=1}^n \alpha_i \omega^i = \sum_{i=1}^n \alpha_i \left(\sum_{k=1}^n s_k^i \hat{\omega}^k \right) = \sum_{k=1}^n \left(\sum_{i=1}^n \alpha_i s_k^i \right) \hat{\omega}^k.$$

Also gilt für das Verhalten der Koordinaten beim Basiswechsel

$$\hat{\alpha}_k = \sum_{i=1}^n \alpha_i s_k^i. \quad (3.1.3)$$

Der Vergleich mit (3.1.1) zeigt, dass sich die Koordinaten von Kovektoren *gleichsinnig* zur Basis in $\mathbb{V}$ transformieren. Man nennt dies auch **kovariantes** Transformationsverhalten (daher auch der Name 'Ko'(varianter)vektor).

Kann man das Spiel weiter treiben und den Dualraum des Dualraums konstruieren? Antwort: Im Prinzip ja, aber es liefert nichts Neues. Man kann zeigen, dass der Doppel-Dualraum $(\mathbb{V}^*)^*$ eines Vektorraums $\mathbb{V}$ kanonisch isomorph zu $\mathbb{V}$ ist.

Übung 3.1: Man zeige, dass jeder Vektor $\mathbf{u} \in \mathbb{V}$ eine lineare Abbildung $\phi_{\mathbf{u}} : \mathbb{V}^* \rightarrow \mathbb{R}$ definiert. Die Abbildung $\mathbb{V} \rightarrow (\mathbb{V}^*)^*$, $\mathbf{u} \mapsto \phi_{\mathbf{u}}$ ist linear und bijektiv, damit ein Isomorphismus.

3.1.2 Lineare Abbildungen

Wir betrachten eine lineare Abbildung $\mathbf{A} : \mathbb{V} \rightarrow \mathbb{V}$. Bezüglich einer Basis lässt sich dieser Abbildung eine **Matrixdarstellung** zuordnen: für jeden Basisvektor $\mathbf{e}_k$ ist $\mathbf{A}\mathbf{e}_k$

²Der Begriff der dualen Basis tritt in der Physik an mehreren Stellen auf; unter anderem auch als *reziprokes Gitter* in der Festkörperphysik

eine Linearkombination

$$\mathbf{A} \mathbf{e}_k = \sum_{i=1}^n A^i_k \mathbf{e}_i.$$

Die n^2 Skalare bilden eine $n \times n$ -Matrix, mit der man nun das Bild eines jeden Vektors berechnen kann, sofern man seine Koordinaten u^k bzgl. der Basis $\mathbf{e}_k$ kennt:

$$\mathbf{A} \mathbf{u} = \sum_{k=1}^n u^k \mathbf{A} \mathbf{e}_k = \sum_{k=1}^n u^k \left(\sum_{i=1}^n A^i_k \mathbf{e}_i \right) = \sum_{i=1}^n \left(\sum_{k=1}^n A^i_k u^k \right) \mathbf{e}_i.$$

D.h. die Koordinaten v^i des Bildvektors $\mathbf{v} = \mathbf{A} \mathbf{u}$ bzgl. einer Basis $\mathbf{e}_k$ sind gegeben durch

$$v^i = \sum_{k=1}^n A^i_k u^k,$$

also durch Matrixmultiplikation der Matrixdarstellung von $\mathbf{A}$ mit den Koordinaten des Vektors bzgl. der Basis $\mathbf{e}_k$. Bei Basiswechsel gilt natürlich

$$\hat{v}^i = \sum_{k=1}^n \hat{A}^i_k \hat{u}^k,$$

d.h., wir haben wieder das gleiche Resultat wie oben: *die Regel, wie die Bilder unter Abbildung zu berechnen sind, ist unabhängig von der Wahl der Basis*. Haben wir also das Bild eines Vektors unter einer Abbildung zu bestimmen, so wählen wir eine Basis, bestimmen die Koordinatendarstellungen von Vektor und Abbildung bzgl. dieser Basis, führen die Matrixmultiplikation durch und erhalten so die Koordinaten des Bildvektors.

3.1.3 Tensoren

Zum Schluss betrachten wir *multilineare Abbildungen* $T : (\mathbb{V}^*)^r \times \mathbb{V}^s \rightarrow \mathbb{R}$, sogenannte Multilinearformen oder **Tensoren**. Ein (r, s) -Tensor ist also eine reellwertige Abbildung mit $r + s$ Argumenten, von denen r bzw. s Stück in $\mathbb{V}^*$ bzw. $\mathbb{V}$ liegen,

Tensoren

$$T(\alpha^1, \dots, \alpha^r, \mathbf{v}_1, \dots, \mathbf{v}_s) \in \mathbb{R}$$

und die in jedem Argument linear ist. Die Anzahl der Argumente nennt man auch **Tensorstufe**. Der Tensor heißt r -fach kontravariant und s -fach kovariant, wenn r bzw. s seiner Argumente aus $\mathbb{V}^*$ bzw. $\mathbb{V}$ sind.

Tensorstufe

Ganz analog zu einer gewöhnlichen linearen Abbildung, lässt sich auch einem (r, s) -Tensor eine Koordinatendarstellung zuordnen. Wir wählen eine Basis $\mathbf{e}_i$ in $\mathbb{V}$, ihre dual-

le Basis ω^i in $\mathbb{V}^*$ und schreiben

$$\begin{aligned}\alpha^1 &= \sum_{i_1=1}^n \alpha_{i_1}^1 \omega^{i_1}, & \mathbf{v}_1 &= \sum_{k_1=1}^n v_1^{k_1} \mathbf{e}_{k_1}, \\ \alpha^2 &= \sum_{i_2=1}^n \alpha_{i_2}^2 \omega^{i_2}, & \mathbf{v}_2 &= \sum_{k_2=1}^n v_2^{k_2} \mathbf{e}_{k_2}, \\ &\vdots & &\vdots \\ \alpha^r &= \sum_{i_r=1}^n \alpha_{i_r}^r \omega^{i_r}, & \mathbf{v}_s &= \sum_{k_s=1}^n v_s^{k_s} \mathbf{e}_{k_s}.\end{aligned}$$

Einsetzen in T und Ausnutzen der Multilinearität ergibt nach etwas Rechnerei

$$\begin{aligned}T(\alpha^1, \dots, \alpha^r, \mathbf{v}_1, \dots, \mathbf{v}_s) \\ = \sum_{\substack{i_1, \dots, i_r=1 \\ k_1, \dots, k_s=1}}^n \alpha_{i_1}^1 \alpha_{i_2}^2 \dots \alpha_{i_r}^r v_1^{k_1} v_2^{k_2} \dots v_s^{k_s} T^{i_1 i_2 \dots i_r}_{k_1 k_2 \dots k_s},\end{aligned}\quad (3.1.4)$$

wobei wir die $n^{(r+s)}$ Skalare $T^{i_1 i_2 \dots i_r}_{k_1 k_2 \dots k_s}$ durch Auswerten der Abbildung T auf allen möglichen Kombinationen von Basisvektoren erhalten

$$T^{i_1 i_2 \dots i_r}_{k_1 k_2 \dots k_s} := T(\omega^{i_1}, \omega^{i_2}, \dots, \omega^{i_r}, \mathbf{e}_{k_1}, \mathbf{e}_{k_2}, \dots, \mathbf{e}_{k_s}).$$

Wir haben (3.1.4) wieder als Rechenvorschrift zu interpretieren: wollen wir das Bild der Abbildung berechnen, dann wählen wir wieder eine Basis in $\mathbb{V}$, die zugehörige duale Basis in $\mathbb{V}^*$, berechnen die Koordinatendarstellung von T und erhalten das Bild schließlich durch Ausführen der länglichen Summen. Wieder ist wichtig festzuhalten, dass diese Rechenvorschrift nicht von der Wahl der Basis abhängt.

Beispiele

- Jeder Kovektor $\alpha \in \mathbb{V}^*$ ist ein $(0, 1)$ -Tensor, also von 1. Stufe und einfach kovariant.
- Ein Vektor $\mathbf{u} \in \mathbb{V}$ ist ein $(1, 0)$ -Tensor also 1. Stufe und einfach kontravariant, denn er definiert eine Abbildung $\mathbb{V}^* \rightarrow \mathbb{R}$, $\alpha \mapsto \langle \alpha, \mathbf{u} \rangle$.

Kronecker-Tensor

- Der $(1, 1)$ -**Kronecker-Tensor** δ ist definiert durch die Abbildung $\mathbb{V}^* \times \mathbb{V} \rightarrow \mathbb{R}$, $(\alpha, \mathbf{u}) \mapsto \langle \alpha, \mathbf{u} \rangle$.

Übung 3.2: Wie sieht die Koordinatendarstellung bzgl. einer Basis aus? Und wie bzgl. einer beliebig anderen Basis?

- Jede lineare Abbildung $\mathbb{V} \rightarrow \mathbb{V}$ lässt sich als Tensor 2. Stufe, einfach ko- und einfach kontravariant auffassen.

Übung 3.3: Wie sieht die entsprechende Multilinearform aus. Was ergibt sich für die Identitätsabbildung?

Übung 3.4: Wie erklärt sich die Bezeichnung r -fach kontravariant und s -fach kovariant für einen (r, s) -Tensor?

Mit Tensoren kann man rechnen. Es gibt folgende Rechenoperationen:

- (i) **Linearkombination** von (r, s) -Tensoren T und S mit Skalaren p und q :

$$\begin{aligned} (pT + qS)(\alpha_1, \dots, \alpha_r, \mathbf{u}_1, \dots, \mathbf{u}_s) \\ := pT(\alpha_1, \dots, \alpha_r, \mathbf{u}_1, \dots, \mathbf{u}_s) + qS(\alpha_1, \dots, \alpha_r, \mathbf{u}_1, \dots, \mathbf{u}_s). \end{aligned}$$

- (ii) **Äußeres oder Tensor-Produkt** von zwei beliebigen Tensoren: ist T_1 bzw. T_2 ein (r_1, s_1) - bzw. (r_2, s_2) -Tensor, dann ist $T_1 \otimes T_2$ ein $(r_1 + r_2, s_1 + s_2)$ -Tensor, definiert durch die Vorschrift

$$\begin{aligned} T_1 \otimes T_2(\alpha_1, \dots, \alpha_{r_1+r_2}, \mathbf{u}_1, \dots, \mathbf{u}_{s_1+s_2}) \\ := T_1(\alpha_1, \dots, \alpha_{r_1}, \mathbf{u}_1, \dots, \mathbf{u}_{s_1}) T_2(\alpha_{r_1+1}, \dots, \alpha_{r_1+r_2}, \mathbf{u}_{s_1+1}, \dots, \mathbf{u}_{s_1+s_2}) \end{aligned}$$

- (iii) **Verjüngung** oder Kontraktion eines (r, s) -Tensors T , $r, s \geq 1$, definiert durch

$$\begin{aligned} \mathcal{C}_j^i T(\alpha_1, \dots, \hat{\alpha}_i, \dots, \alpha_r, \mathbf{u}_1, \dots, \hat{\mathbf{u}}_j, \dots, \mathbf{u}_s) \\ := \sum_{k=1}^n T(\alpha_1, \dots, \omega^k, \dots, \alpha_r, \mathbf{u}_1, \dots, \mathbf{e}_k, \dots, \mathbf{u}_s) \end{aligned}$$

d.h. die Verjüngung über die i, j Stellen entsteht dadurch, dass man als j -tes Argument einen Basisvektor und als i -tes Argument den entsprechenden dualen Basisvektor verwendet und dann über alle Basisvektoren summiert.

- (iv) Schließlich gibt es noch die Operation der **Argumentvertauschung**: Ist T ein (r, s) -Tensor, dann ist durch

$$T'(\alpha_1, \dots, \alpha_i, \dots, \alpha_j, \dots, \mathbf{u}_1, \dots) := T(\alpha_1, \dots, \alpha_j, \dots, \alpha_i, \dots, \mathbf{u}_1, \dots)$$

ein weiterer (r, s) -Tensor definiert, der aus T durch Vertauschen der Argumente α_i und α_j entsteht. Entsprechendes gilt für das Vertauschen der Argumente aus $\mathbb{V}$.

Es ist üblich die Skalare, also den Körper über dem $\mathbb{V}$ definiert ist (hier also $\mathbb{R}$) mit den $(0, 0)$ -Tensoren gleichzusetzen. Die (r, s) -Tensoren bilden unter (i) einen Vektorraum.

Übung 3.5: Welche Dimension besitzt dieser Vektorraum?

Die Menge aller Tensoren bildet unter (i) und (ii) eine Algebra, die **Tensoralgebra** über $\mathbb{V}$. Diese Tensoralgebra kann über jedem Vektorraum konstruiert werden.

Es ist offensichtlich, dass die Struktur eines Tensors etwas kompliziert ist. Die Angabe eines Symbols, z.B. T alleine, genügt nicht, um die Stufe oder die Anzahl von ko- und kontravarianten Argumenten anzugeben, geschweige denn, um damit rechnen zu

können. Andererseits ist das explizite Rechnen mit diesen Größen ebenfalls aufwendig. Auch wenn das Ergebnis basisunabhängig ist, muß man bei solchen Rechnungen immer eine Basis spezifizieren, bzgl. der man die Rechnungen ausführt. Wir werden daher eine Notation verwenden, die beiden Bedürfnissen Rechnung trägt: expliziter als die rein mathematische Notation und nicht ganz so explizit wie beim konkreten Rechnen.

Index-Notation Wir stellen einen (r, s) -Tensor als ein Symbol

$$T^{a\dots b}_{c\dots d}$$

mit r oberen und s unteren Indizes. Die Indizes sind vom vorderen Teil des Alphabets genommen ($a, b, \dots$). Die Bezeichnung der Indizes ist irrelevant. Wichtig ist nur ihre Stellung relativ zu den anderen Indizes. So bezeichnet z.B. T^{ab}_c einen Tensor gleichen Typs wie T^{de}_f .

Ein Vektor wird als v^a dargestellt, ein Kovektor als α_a . Die duale Paarung $\langle \alpha, u \rangle$ geht über in die Darstellung $\alpha_a u^a$ und ist so zu interpretieren: wenn wir eine Basis und duale Basis wählten, α und u bzgl. diesen Basen darstellten, dann müssten wir die duale Paarung über

$$\sum_{i=1}^n \alpha_i u^i$$

berechnen. Die Index-Notation $\alpha_a u^a$ ist also eine symbolische Erinnerung daran, wie die Ausdrücke im konkreten Fall zu berechnen wären. Man beachte, dass die Summenzeichen weggelassen werden. Dies ist nicht tragisch, wenn man verabredet, dass über gleiche Indizes, die einmal oben und einmal unten erscheinen, summiert werden soll. Dies ist die **Einsteinsche Summenkonvention**. Die Indizes haben keine 'konkrete' Bedeutung in dem Sinne, dass sie als Platzhalter stünden für die Zahlen $1 : n$. Sie haben vielmehr eine 'buchhalterische Bedeutung', denn ihre Stellung deklariert den Typ des Objekts. Sie haben eine ähnliche Funktion wie der Pfeil in $\vec{u}$ oder die Druckerschwärze in u , nur dass sie noch mehr Information tragen.

Einsteinsche
Summenkonvention

Eine lineare Abbildung $A : \mathbb{V} \rightarrow \mathbb{V}$, also einen $(1, 1)$ -Tensor, stellen wir in der Form

$$A^a_b$$

dar, das Bild eines Vektors u^a unter dieser Abbildung wird dann

$$A^a_b u^b. \quad (3.1.5)$$

Insbesondere hat die Identitäts-Abbildung, also der Kronecker-Tensor, die Darstellung δ^a_b und es gilt für jeden Vektor u^a

$$\delta^a_b u^b = u^a.$$

Schließlich betrachten wir noch die Rechenoperationen für Tensoren mithilfe dieser Notation. Die Linearkombination zweier (r, s) -Tensoren $T^{a\dots b}_{c\dots d}$ und $U^{a\dots b}_{c\dots d}$ wird

$$pT^{a\dots b}_{c\dots d} + qU^{a\dots b}_{c\dots d}.$$

Das äußere Produkt zweier beliebiger Tensoren wird

$$\begin{aligned} (T^{a_1 \dots a_r}_{b_1 \dots b_s}, U^{a_1 \dots a_t}_{b_1 \dots b_u}) \\ \mapsto S^{a_1 \dots a_{r+t}}_{b_1 \dots b_{s+u}} = T^{a_1 \dots a_r}_{b_1 \dots b_s} U^{a_1 \dots a_t}_{b_1 \dots b_u}. \end{aligned}$$

So ist z.B. das äußere Produkt zweier Vektoren u^a und v^b gegeben durch den $(2, 0)$ -Tensor

$$T^{ab} = u^a v^b.$$

Die Kontraktion schreibt sich in dieser Notation extrem einfach

$$\mathcal{C}_j^i : T^{a\dots b}_{c\dots d} \mapsto T^{a\dots e\dots b}_{c\dots e\dots d}$$

wobei der Index e an der i -ten Stelle oben und der j -ten Stelle unten steht. Ein Beispiel für eine solche Kontraktion ist die **Spurbildung** bei einer linearen Abbildung

Spurbildung

$$\text{tr} : A^a_b \mapsto A^a_a.$$

Insbesondere gilt für die Identität $\delta^a_a = n = \dim_{\mathbb{R}} \mathbb{V}$. Eine weitere Rechenoperation, die sich durch Kombination von äußerem Produkt und Kontraktion erreichen lässt, ist die **Überschiebung** oder Transvektion eines Tensors $T^{a\dots b}_{c\dots d}$ mit einem Vektor u^a , ausgedrückt in Index-Notation (vgl. (3.1.5))

Überschiebung

$$(T^{a\dots b}_{c\dots d}, u^a) \mapsto T^{a\dots b}_{c\dots e\dots d} u^e.$$

Die Auswertung eines Tensors $T^{a\dots b}_{c\dots d}$ auf seinen Argumenten wird allgemein so dargestellt (vgl. (3.1.4))

$$T^{a\dots b}_{c\dots d} \alpha_a^1 \dots \alpha_b^r u_1^c \dots u_s^d.$$

Damit erklärt sich auch die Operation der Argumentvertauschung in Index-Notation, z.B.

$$T'^{a\dots b\dots c\dots}_{e\dots f} = T^{a\dots c\dots b\dots}_{e\dots f}$$

und entsprechend für die unteren Indizes. Beispielsweise haben wir

$$T'(\mathbf{u}, \mathbf{v}) = T'_{ab} u^a v^b = T_{ba} u^a v^b = T(\mathbf{v}, \mathbf{u}).$$

Da wir logischerweise nur Argumente gleichen Typs also Vektoren bzw. Kovektoren vertauschen können, können wir auch nur obere Indizes mit oberen vertauschen und untere Indizes mit unteren.

Zwei wichtige Operationen sind die **Symmetrisierung** und die **Antisymmetrisierung**

Symmetrisierung

Antisymmetrisierung

eines Tensors. Als Beispiel betrachten wir einen $(0, 2)$ -Tensor T_{ab} . Wir können immer schreiben

$$T_{ab} = \frac{1}{2} (T_{ab} + T_{ba}) + \frac{1}{2} (T_{ab} - T_{ba}) = S_{ab} + A_{ab}.$$

Die beiden Summanden S_{ab} und A_{ab} haben *Symmetrie-Eigenschaften*

$$S_{ab} = S_{ba}, \quad A_{ab} = -A_{ba}.$$

symmetrischer
Tensor
antisymmetrischer
Tensor

Der Tensor S_{ab} ist ein **symmetrischer Tensor**, während A_{ab} ein **antisymmetrischer Tensor** ist. Sie entstehen aus T_{ab} durch Symmetrisierung bzw. Antisymmetrisierung. I.a. wird ein $(0, s)$ -Tensor symmetrisiert durch die Vorschrift

$$T_{a_1 \dots a_s} \mapsto T_{(a_1 \dots a_s)} = \frac{1}{s!} \sum_{\pi} T_{a_{\pi(1)} \dots a_{\pi(s)}}.$$

Dabei ist π eine Permutation der Zahlen $1, 2, \dots, s$. Die Antisymmetrisierung eines Tensors wird durch

$$T_{a_1 \dots a_s} \mapsto T_{[a_1 \dots a_s]} = \frac{1}{s!} \sum_{\pi} (-1)^{\pi} T_{a_{\pi(1)} \dots a_{\pi(s)}},$$

erreicht. Dabei ist $(-1)^{\pi}$ das 'Signum' der Permutation π , also

$$(-1)^{\pi} = \begin{cases} +1 & \pi \text{ ist eine gerade Permutation} \\ -1 & \pi \text{ ist eine ungerade Permutation} \end{cases}.$$

Konkret ist z. B. die (Anti-) Symmetrisierung eines $(0, 3)$ -Tensors gegeben durch

$$V_{(abc)} = \frac{1}{6} (V_{abc} + V_{bca} + V_{cab} + V_{bac} + V_{cba} + V_{acb})$$

$$V_{[abc]} = \frac{1}{6} (V_{abc} + V_{bca} + V_{cab} - V_{bac} - V_{cba} - V_{acb}).$$

total
(anti)symmetrisch

Ein Tensor $T_{a \dots b}$ heißt **total (anti)symmetrisch** falls $T_{a \dots b} = T_{(a \dots b)}$ bzw. $T_{a \dots b} = T_{[a \dots b]}$ gilt. Alle oben genannten Operationen lassen sich natürlich auch für obere Indizes durchführen.

Symmetrie-Eigenschaften von Tensoren sind wichtige Hilfsmittel, die beim Rechnen mit Tensoren oft große Vereinfachungen mit sich bringen. Als Beispiel seien genannt:

1. (Anti-)Symmetrisierung ist ein *Projektor*, d.h. zweimalige Anwendung ändert nichts am Ergebnis. Ist z.B. $T_{abcde} = U_{[abcde]}$, dann ist

$$T_{[abcde]} = U_{[[abcde]]} = U_{[abcde]} = T_{abcde}.$$

2. Andererseits schließen sich Symmetrisierung und Anti-Symmetrisierung gegenseitig aus. Im obigen Beispiel gilt

$$T_{(abcde)} = U_{([abcde])} = 0.$$

Symmetrisierung und anschließende Anti-Symmetrisierung (und umgekehrt) ergeben Null.

Übung 3.6: Zeigen Sie diese Eigenschaft.

3. Für einen symmetrischen $(0, 2)$ -Tensor S_{ab} und einen antisymmetrischen $(2, 0)$ -Tensor A^{ab} gilt identisch

$$A^{ab}S_{ab} = 0.$$

Übung 3.7: Zeigen Sie diese Eigenschaft.

4. Für einen Tensor V_{abc} mit $V_{abc} = V_{a[bc]}$ und $V_{abc} = V_{(ab)c}$ gilt identisch

$$V_{abc} = 0.$$

Übung 3.8: Zeigen Sie diese Eigenschaft.

3.2 Anwendung: Euklidische Geometrie, Vektorrechnung

Als kurze Anwendung des entwickelten Tensor-Kalküls betrachten wir nun die Euklidische Geometrie, ohne jedoch zu sehr in die Tiefe zu gehen. Das Ziel ist es, bekannte geometrische Strukturen mit den Augen eines 'Tensorikers' zu betrachten und erste Rechnungen zu machen.

Wir betrachten wie bisher einen (endlich dimensional) Vektorraum $\mathbb{V}$ über $\mathbb{R}$, der nun aber zusätzlich mit einem inneren Produkt, auch oft **Metrik** genannt, ausgestattet sei, also einen Hilbertraum. Das innere Produkt ist eine Abbildung

$$\mathbb{V} \times \mathbb{V} \rightarrow \mathbb{R}, (\mathbf{u}, \mathbf{v}) \mapsto \langle \mathbf{u} | \mathbf{v} \rangle$$

mit den Eigenschaften

1. $\langle \cdot | \cdot \rangle$ ist *bilinear*.
2. $\langle \cdot | \cdot \rangle$ ist *symmetrisch*: für alle $\mathbf{u}, \mathbf{v} \in \mathbb{V}$ ist

$$\langle \mathbf{u} | \mathbf{v} \rangle = \langle \mathbf{v} | \mathbf{u} \rangle.$$

3. $\langle \cdot | \cdot \rangle$ ist *positiv definit*:

$$\langle \mathbf{u} | \mathbf{u} \rangle > 0, \quad \text{für alle } \mathbf{u} \neq 0.$$

Die hier wesentliche Eigenschaft ist zunächst die erste. Das innere Produkt ist eine *Bilinearform* auf $\mathbb{V}$ und damit natürlich ein $(0, 2)$ -Tensor, den wir mit

$$g_{ab}$$

bezeichnen. Die zweite Eigenschaft charakterisiert diesen Tensor als symmetrisch, also

$$g_{ab} = g_{ba} = g_{(ab)} \iff g_{ab}u^a v^b = g_{ab}v^a u^b.$$

Die dritte Eigenschaft schließlich schreibt sich in Index-Notation

$$g_{ab}u^a u^b > 0, \quad \text{für alle } u^a \neq 0.$$

Wesentlich an dieser Eigenschaft ist für uns nicht die Positivität, sondern die daraus resultierende *Regularität* der Bilinearform:

$$\langle \mathbf{u} | \cdot \rangle = 0 \iff \mathbf{u} = 0.$$

Diese Regularität erlaubt es nämlich $\mathbb{V}$ mit $\mathbb{V}^*$ bzw. Vektoren und Kovektoren zu identifizieren. Jeder Vektor $\mathbf{u} \in \mathbb{V}$ definiert eine Linearform, also eine lineare Abbildung $\mathbb{V} \rightarrow \mathbb{R}$ via

$$\mathbf{v} \mapsto \langle \mathbf{u} | \mathbf{v} \rangle.$$

Die so vermittelte Abbildung $\mathbb{V} \rightarrow \mathbb{V}^*$, $\mathbf{u} \mapsto \langle \mathbf{u} | \cdot \rangle$ ist injektiv und daher bijektiv. Sie ist also ein kanonischer Isomorphismus, der es erlaubt Vektoren und Kovektoren, die ein-eindeutig aufeinander abgebildet werden, zu identifizieren.

Bemerkung *Ist auf einem (endlich dimensionalen) Vektorraum $\mathbb{V}$ eine reguläre Bilinearform definiert, dann sind $\mathbb{V}$ und $\mathbb{V}^*$ isomorph. Man muss daher nicht zwischen Vektoren und Kovektoren unterscheiden.*

In Index-Notation stellt sich dieser Sachverhalt so dar. Die Bilinearform g_{ab} definiert die Abbildung $\mathbb{V} \rightarrow \mathbb{V}^*$

$$u^a \mapsto g_{ab}u^b =: U_a$$

Die Index-Stellung an U_a signalisiert automatisch, dass es sich hierbei um einen Kovektor handelt. Die Tatsache, dass die Abbildung $u^a \mapsto U_a$ bijektiv ist, erlaubt es uns, das Bild von u^a unter dieser Abbildung mit u_a zu bezeichnen, denn es kann keine Mehrdeutigkeit geben. Wir haben also

$$g_{ab}u^a = u_b,$$

Senken von Indizes

was zu der Rechenregel Anlaß gibt: Die Metrik wird zum **Senken von Indizes** verwendet. Die Umkehrabbildung $\mathbb{V}^* \rightarrow \mathbb{V}$ ordnet jedem Kovektor α_b den eindeutig definierten Vektor α^a zu, der dadurch definiert ist, dass für alle Vektoren v^b

$$\alpha_b v^b = g_{ab} \alpha^a v^b$$

gilt. Sie hat also die Index-Notation G^{ab} und es gilt

$$u^a = G^{ab} g_{bc} u^c = \delta_c^a u^c$$

für beliebige Vektoren u^a . Es ist also $G^{ab} g_{bc} = \delta_c^a$ und wir erhalten damit die weitere Beziehung

$$G^{ab} g_{ac} g_{bd} = g_{ac} \delta_d^a = g_{cd},$$

travariante
Metrik

d.h., werden die Indizes von G^{ab} mit der Metrik gesenkt, erhalten wir gerade die Metrik zurück. Man nennt daher oft G^{ab} die **kontravariante Metrik** und bezeichnet sie mit g^{ab} .

von Indizes

Übung 3.9: Zeigen Sie, dass $g^{ab} = g^{ba}$ gilt.

Die kontravariante Metrik kann zum **Heben von Indizes** verwendet werden.

Wollen wir mit der Metrik explizit rechnen, müssen wir eine Basis einführen. Die Metrik erlaubt es uns, eine besondere Klasse von Basen zu betrachten, die Orthonormalbasen. Eine **Orthonormalbasis** (ONB) ist ein n -tupel von Vektoren $(e_1^a, e_2^a, \dots, e_n^a)$ mit der Eigenschaft

Orthonormalbasis

$$g_{ab}e_i^a e_k^b = \begin{cases} 1 & \text{für } i = k \\ 0 & \text{für } i \neq k \end{cases}$$

Die Matrix- oder Koordinatendarstellung des Tensors g_{ab} bzgl. einer ONB ist daher

$$(g_{ik})_{i,k=1:n} = \begin{pmatrix} 1 & 0 & 0 & \dots \\ 0 & 1 & 0 & \dots \\ 0 & 0 & 1 & \dots \\ \vdots & \vdots & & \ddots \end{pmatrix},$$

die Einheitsmatrix. Bezüglich einer beliebigen Basis ist dies offensichtlich nicht der Fall. Die Matrix ist jedoch immer notwendigerweise symmetrisch und regulär (ihre Determinante ist von Null verschieden).

Soweit zur Metrik. Es gibt noch einen weiteren Tensor, der auf einem n -dimensionalen Vektorraum $\mathbb{V}$ definiert werden kann. Je n Vektoren $u_1, \dots, u_n$ spannen ein Parallelepiped auf. Dessen Volumen $\mathcal{V}$ ist eine reelle Zahl. Die Abbildung, die jedem n -tupel von Vektoren, das Volumen des von ihnen aufgespannten Parallelepipeds zuordnet, ist linear. Es handelt sich also wieder um einen Tensor, und zwar n -fach kovariant. Er heißt **Volumenform** und wir bezeichnen ihn mit

Volumenform

$$\varepsilon_{ab\dots c},$$

Damit ist $\mathcal{V} = \varepsilon_{ab\dots c} u_1^a u_2^b \dots u_n^c$. Sind in diesem Ausdruck zwei der Vektoren einander gleich, dann ist das aufgespannte Parallelepiped entartet, sein Volumen $\mathcal{V} = 0$. Es ist also

$$\varepsilon_{a\dots b\dots c\dots d} u_1^a \dots u_i^b \dots u_k^c \dots u_n^d = 0$$

falls für $i \neq k$ die Vektoren $u_i^a = u_k^a$ gleich sind.

Übung 3.10: Zeigen Sie, dass der Tensor $\varepsilon_{ab\dots c}$ total antisymmetrisch ist

$$\varepsilon_{ab\dots c} = \varepsilon_{[ab\dots c]}.$$

Man kann zeigen, dass je zwei total antisymmetrische $(0, n)$ -Tensoren auf einem n -dimensionalen Vektorraum einander proportional sein müssen. Um die Volumenform eindeutig festlegen zu können, müssen wir daher nur angeben, welchen Wert sie für

eine Basis annimmt. Wir setzen also das Volumen eines n -dimensionalen Einheitswürfels, also desjenigen Parallelepipeds, das von einer ONB aufgespannt wird gleich 1. Es gilt also für eine ONB

$$\varepsilon_{ab\dots c} e_1^a e_2^b \dots e_n^c = \varepsilon_{12\dots n} = 1,$$

und daher

$$\varepsilon_{i_1 i_2 \dots i_n} = \begin{cases} +1 & \text{falls } (i_1, i_2, \dots, i_n) \text{ eine gerade Permutation ist,} \\ -1 & \text{falls } (i_1, i_2, \dots, i_n) \text{ eine ungerade Permutation ist,} \\ 0 & \text{sonst} \end{cases}$$

Übung 3.11: Was ergibt sich für eine beliebige Basis?

Zum Schluß spezialisieren wir uns auf den 3-dimensionalen Fall. Wir haben dann eine Metrik g_{ab} und einen total antisymmetrischen $(0,3)$ -Tensor ε_{abc} vorliegen. Sei (e_1^a, e_2^a, e_3^a) eine ONB und $(\omega_a^1, \omega_a^2, \omega_a^3)$ die dazu duale Basis (warum ist $\omega_a^i = g_{ab} e_i^b$ für $i = 1, 2, 3$?) Aus dieser dualen Basis konstruieren wir einen total antisymmetrischen $(0,3)$ -Tensor

$$\Omega_{abc} := \omega_{[a}^1 \omega_b^2 \omega_{c]}^3.$$

Da solche Tensoren proportional zueinander sind, muß es eine reelle Zahl c geben, so dass $\Omega_{abc} = c \varepsilon_{abc}$ gilt. Wir bestimmen diese Zahl, indem wir die beiden Formen auf einer Basis auswerten. Dazu verwenden wir natürlich am geschicktesten die vorgegebene ONB. Es ist

$$\begin{aligned} \Omega_{abc} = \frac{1}{6} & \left(\omega_a^1 \omega_b^2 \omega_c^3 + \omega_a^2 \omega_b^3 \omega_c^1 + \omega_a^3 \omega_b^1 \omega_c^2 \right. \\ & \left. - \omega_a^2 \omega_b^1 \omega_c^3 - \omega_a^3 \omega_b^2 \omega_c^1 - \omega_a^1 \omega_b^3 \omega_c^2 \right), \end{aligned}$$

so dass

$$\Omega_{abc} e_1^a e_2^b e_3^c = \frac{1}{6}$$

und wegen $\varepsilon_{abc} e_1^a e_2^b e_3^c = 1$ folgt $c = \frac{1}{6}$ und damit

$$\varepsilon_{abc} = 6 \omega_{[a}^1 \omega_b^2 \omega_{c]}^3.$$

Wir können nun auch einen total antisymmetrischen $(3,0)$ -Tensor ε^{abc} einführen, indem wir die Indizes an ε_{abc} mit der Metrik heben. Mit einer ganz analogen Überlegung wie oben bekommen wir das Resultat

$$\varepsilon^{abc} = 6 e_1^{[a} e_2^b e_3^{c]}.$$

Wir interessieren uns jetzt für den $(3,3)$ -Tensor

$$\varepsilon_{abc} \varepsilon^{def}.$$

Dieser muss proportional sein zu dem (3, 3)-Tensor

$$\delta_a^{[d} \delta_b^e \delta_c^{f]},$$

denn beide Tensoren sind total antisymmetrisch jeweils in ihren oberen und unteren Indizes. Es gibt also wieder eine reelle Zahl c , so dass

$$\varepsilon_{abc} \varepsilon^{def} = c \delta_a^{[d} \delta_b^e \delta_c^{f]}$$

gilt. Um die Zahl c zu bestimmen, benutzen wir wieder die ONB. Da $\delta_a^d e_i^a = e_i^d$ für $i = 1, 2, 3$ gilt erhalten wir

$$\varepsilon^{def} = \varepsilon^{def} \varepsilon_{abc} e_1^a e_2^b e_3^c = c \delta_a^{[d} \delta_b^e \delta_c^{f]} e_1^a e_2^b e_3^c = c e_1^{[d} e_2^e e_3^{f]} = \frac{c}{6} \varepsilon^{def},$$

also ist $c = 6$. Wozu ist die Identität

$$\varepsilon_{abc} \varepsilon^{def} = 6 \delta_a^{[d} \delta_b^e \delta_c^{f]}$$

nun gut? Zunächst einmal erhalten wir daraus noch einmal drei weitere Identitäten durch sukzessive Kontraktion über die Indizes

$$\varepsilon_{abc} \varepsilon^{aef} = \delta_b^e \delta_c^f - \delta_b^f \delta_c^e, \quad (3.2.1)$$

$$\varepsilon_{abc} \varepsilon^{abf} = 2 \delta_c^f, \quad (3.2.2)$$

$$\varepsilon_{abc} \varepsilon^{abc} = 6. \quad (3.2.3)$$

Übung 3.12: Leiten Sie diese Relationen her.

Es seien zwei Vektoren u^a, v^a gegeben, dann können wir den Kovektor

$$\varepsilon_{abc} u^b v^c$$

bilden und danach durch Heben des Index einen Vektor w^a konstruieren

$$w^a = g^{ad} \varepsilon_{abc} u^b v^c.$$

Damit haben wir je zwei Vektoren auf antisymmetrische bilineare Weise einen weiteren Vektor zugeordnet. Es handelt sich hierbei also um ein Vektorprodukt.

Übung 3.13: Überzeugen Sie sich davon, dass dies genau das übliche **Kreuzprodukt** ist,

Kreuzprodukt

$$\mathbf{w} = \mathbf{u} \times \mathbf{v}.$$

Im Übrigen ist für beliebige Vektoren u^a, v^b, w^c

$$\varepsilon_{abc} u^a v^b w^c = \langle \mathbf{w} | \mathbf{u} \times \mathbf{v} \rangle.$$

Bekanntlich gibt es zwischen innerem Produkt und Kreuzprodukt einige Identitäten:

$$\mathbf{u} \times (\mathbf{v} \times \mathbf{w}) = \mathbf{v} \langle \mathbf{u} | \mathbf{w} \rangle - \mathbf{w} \langle \mathbf{u} | \mathbf{v} \rangle \quad (3.2.4)$$

$$\langle \mathbf{u} \times \mathbf{v} | \mathbf{w} \times \mathbf{t} \rangle = \langle \mathbf{u} | \mathbf{w} \rangle \langle \mathbf{v} | \mathbf{t} \rangle - \langle \mathbf{u} | \mathbf{t} \rangle \langle \mathbf{v} | \mathbf{w} \rangle, \quad (3.2.5)$$

$$(\mathbf{u} \times \mathbf{v}) \times (\mathbf{w} \times \mathbf{t}) = \langle \mathbf{u} | \mathbf{w} \times \mathbf{t} \rangle \mathbf{v} - \langle \mathbf{v} | \mathbf{w} \times \mathbf{t} \rangle \mathbf{u} \quad (3.2.6)$$

Um (3.2.6) zu beweisen, formulieren wir die linke Seite in Index-Notation

$$\varepsilon^a{}_{bc} \left(\varepsilon^b{}_{de} u^d v^e \right) \left(\varepsilon^c{}_{fh} w^f t^h \right) = u^d v^e w^f t^h \left(\varepsilon^a{}_{bc} \varepsilon^b{}_{de} \varepsilon^c{}_{fh} \right)$$

und verarbeiten das dreifache Produkt der ε 's weiter. Wir erhalten mit Hilfe von (3.2.1)

$$\begin{aligned} \varepsilon^a{}_{bc} \varepsilon^b{}_{de} \varepsilon^c{}_{fh} &= \left(\varepsilon^a{}_{bc} \varepsilon^b{}_{de} \right) \varepsilon^c{}_{fh} = \left(\varepsilon^{abc} \varepsilon_{bde} \right) \varepsilon_{cfh} = \left(\varepsilon^{bca} \varepsilon_{bde} \right) \varepsilon_{cfh} \\ &= (\delta_d^c \delta_e^a - \delta_d^a \delta_e^c) \varepsilon_{cfh} = \varepsilon_{dfh} \delta_e^a - \varepsilon_{efh} \delta_d^a \end{aligned}$$

Setzen wir dies wieder oben ein, so ergibt sich

$$u^d v^e w^f t^h (\varepsilon_{dfh} \delta_e^a - \varepsilon_{efh} \delta_d^a) = \left(\varepsilon_{dfh} u^d w^f t^h \right) v^a - \left(\varepsilon_{efh} v^e w^f t^h \right) u^a$$

und damit die rechte Seite von (3.2.6).

Übung 3.14: Beweisen Sie die übrigen Identitäten.

4 Lorentz-Geometrie und Minkowski-Raum

In diesem Kapitel wird ein kurzer Abriss der Speziellen Relativitätstheorie gegeben. Das Ziel ist es, den geometrischen Rahmen zu präsentieren, in dem sich die lokale Physik abspielt, d.h., wenn man gravitative Wechselwirkungen ignoriert. Da dies im Hinblick auf die spätere allgemeine Theorie geschieht, ist der Zugang *geometrisch*, die physikalische Motivation tritt in den Hintergrund. Zu diesem Aspekt sei auf die einschlägige Literatur verwiesen¹.

4.1 Die relativistische Raumzeit

Gemäß dem Galileischen **Relativitätsprinzip** läuft die Physik in jedem Inertialsystem gleich ab. Hat man einmal verabredet, was man unter dem Begriff 'Inertialsystem' verstehen will, dann folgt aus dem Relativitätsprinzip, wie sich die Raumzeit mathematisch beschreiben lässt². Beobachten wir z.B. in einem Inertialsystem eine gleichförmige Bewegung, dann ist diese in jedem Inertialsystem gleichförmig. Eine solche Bewegung wird durch eine lineare Beziehung zwischen den inertialen Raum- und Zeitkoordinaten charakterisiert. Beim Übergang von einem Inertialsystem zum anderen muß diese Eigenschaft erhalten bleiben. Daher kann es sich bei der Transformation zwischen zwei Inertialsystemen nur um eine **affine Abbildung** handeln. Dies bedeutet:

Relativitätsprinzip

affine Abbildung

*Die speziell-relativistische Raumzeit ist ein **affiner Raum**.*

affiner Raum

Darunter ist folgendes zu verstehen. Die Raumzeit $\mathbb{M}$ ist eine Menge von **Ereignissen**, Raumzeitpunkten wie 'hier-jetzt', die durch die Angabe von vier reellen Zahlen $(\xi^0, \xi^1, \xi^2, \xi^3)$, den Koordinaten, charakterisiert werden können. Die affine Struktur entsteht durch die Angabe einer Abbildung $\mathbb{M} \times \mathbb{M} \rightarrow \mathbb{V}$ in einen Vektorraum $\mathbb{V}$, die je zwei Ereignissen $P, Q \in \mathbb{M}$ einen Vektor $\text{vec}(P, Q)$ zuordnet und die Eigenschaft besitzt, dass für je drei Ereignisse $O, P, Q \in \mathbb{M}$

Ereignisse

$$\text{vec}(O, P) + \text{vec}(P, Q) = \text{vec}(O, Q)$$

¹vgl. auch das Skript zur Vorlesung 'Spezielle Relativitätstheorie', erhältlich über die URL <http://www.tat.physik.uni-tuebingen.de/~joergf>

²Zur logischen Herleitung der mathematischen Struktur der Raumzeit sei das Buch von H. Weyl *Raum-Zeit-Materie* wärmstens empfohlen.

Translation gilt. Dann gilt $\text{vec}(P, P) = 0$ und $\text{vec}(P, Q) = -\text{vec}(Q, P)$. Wir können uns den Vektor $\text{vec}(PQ)$ als **Translation** vorstellen, die das Ereignis P in das Ereignis Q verschiebt (räumlich und/oder zeitlich). Wählen wir uns ein Ereignis O , den **Ursprung** ein für allemal aus, dann bekommen wir eine Abbildung, die jedem Ereignis P den Vektor $\text{vec}(OP)$ zuordnet und wir fordern, dass diese Abbildung bijektiv sei. Der Vektor $\text{vec}(OP)$ Ortsvektor wird dann als **Ortsvektor** von P bzgl. des Ursprungs O bezeichnet.

Ereignisse in der Raumzeit lassen sich demnach *eindeutig* charakterisieren, indem man ihren Ortsvektor bzgl. eines Ursprungs angibt. Man beachte jedoch, dass diese Beschreibung von der Wahl des Ursprungs abhängt. Ändert man den Ursprung, so ändert sich auch der Ortsvektor. Gewissermaßen haben wir für jede Wahl eines Ursprungs-Ereignisses den Raum der Ortsvektoren bzgl. dieses Ursprungs zu betrachten. Die Struktur der Raumzeit der speziellen Relativitätstheorie ist gerade so, dass diese 'lokalen' Vektorräume identisch sind. Es sind gerade die Translationen, die die lokalen Vektorräume aufeinander abbilden.

Das zweite von Einstein aufgestellte Prinzip formuliert die 'Konstanz der Lichtgeschwindigkeit': *Die Ausbreitungsgeschwindigkeit des Lichts im Vakuum ist unabhängig vom Bewegungszustand der Quelle.* Aus diesem Prinzip ergibt sich eine geometrische Konsequenz für den Vektorraum $\mathbb{V}$:

Lorentz-Metrik *Der Vektorraum $\mathbb{V}$ trägt eine **Lorentz-Metrik**.*

Minkowski-Raumzeit Dadurch wird dem der Raumzeit zugeordneten Vektorraum und damit der Raumzeit selber eine wesentliche Struktur aufgeprägt. Man nennt die so definierte Raumzeit die **Minkowski-Raumzeit** $\mathbb{M}$ und den ihr zugeordneten Vektorraum, den Minkowski-Vektorraum.

Es gibt also auf $\mathbb{V}$ einen symmetrischen $(0, 2)$ -Tensor η_{ab} mit der Eigenschaft, dass seine Matrixdarstellung $(\eta_{ik})_{i,k=0,3}$ bzgl. einer beliebigen Basis, welche immer eine symmetrische Matrix ist, auf die Normalform

$$\eta = \begin{pmatrix} 1 & & & \\ & -1 & & \\ & & -1 & \\ & & & -1 \end{pmatrix}$$

Inertialsystem gebracht werden kann. Die *Lorentz-Signatur* $(1, -1, -1, -1)$ ist charakteristisch für die Metrik, in dem Sinne, dass sie *unabhängig von der gewählten Basis ist* (Sylvestersches Trägheitstheorem). Eine Basis, bzgl. der die Matrixdarstellung genau diese Normalform besitzt, heißt (Pseudo-) ONB, (lokales) **Inertialsystem** oder (lokales) Bezugssystem.

Die so definierte Bilinearform ist regulär und wir können daher $\mathbb{V}$ und $\mathbb{V}^*$ identifizieren, also in der Index-Schreibweise Indizes mit der Metrik senken und mit der kontravarianten Metrik heben.

Die unterschiedlichen Vorzeichen in der Signatur haben schwerwiegende Konsequenzen: es gibt unterschiedliche Klassen von Vektoren: wir können für jeden Vektor u^a sein **Betragsquadrat** $\eta_{ab}u^au^b$ bilden. Nun ist dieses Betragsquadrat kein echtes Quadrat, denn diese Zahl kann positiv wie negativ sein. Wir nennen daher einen beliebigen Vektor v^a **zeitartig, raumartig, bzw. lichtartig**, wenn $\eta_{ab}v^av^b$ positiv, negativ oder Null ist. Eine weitere Unterteilung ergibt sich für **kausale Vektoren**, also zeit- und lichtartige Vektoren. Wir wählen einen beliebigen zeitartigen Vektor t^a fest aus und nennen jeden kausalen Vektor u^a **zukunftsweisend**, wenn $\eta_{ab}t^au^b > 0$ gilt. Entsprechend definieren wir als **vergangenheitsweisend** jene kausalen Vektoren für die $\eta_{ab}t^au^b < 0$. Dadurch haben wir eine **Zeit-Orientierung** für kausale Vektoren festgelegt.

Betragsquadrat

kausale Vektoren

zukunftsweisend
vergangenheitsweisend

Zeit-Orientierung

Übung 4.1: Zeigen Sie, dass diese Zeit-Orientierung unabhängig von der Wahl des Referenzvektors t^a ist in dem Sinne, dass die Zeit-Orientierung bzgl. eines anderen zeitartigen Vektors, der bzgl. t^a zukunftsgerichtet ist, mit der Zeit-Orientierung bzgl. t^a übereinstimmt.

Übung 4.2: Zeigen Sie, dass Vektoren, für die $\eta_{ab}t^au^b = 0$ ist, notwendigerweise raumartig sind.

Es gibt also fünf unterschiedliche Klassen von Vektoren. Wie teilen sich diese den ganzen Vektorraum auf? Zur Beantwortung dieser Frage wählen wir eine ONB und berechnen das Skalarprodukt eines Vektors $v^a = v^0e_0^a + v^1e_1^a + v^2e_2^a + v^3e_3^a$. Es folgt

$$\eta_{ab}v^av^b = (v^0)^2 - (v^1)^2 - (v^2)^2 - (v^3)^2$$

Für lichtartige Vektoren gilt dann

$$v^0 = \pm \sqrt{(v^1)^2 + (v^2)^2 + (v^3)^2}$$

Diese Vektoren bilden also einen (Doppel-) Kegel, den **Nullkegel**. Der erste Basisvektor e_0^a ist zeitartig, so dass wir ihn zur Festlegung der Zeit-Orientierung benutzen können. Dann sind kausale Vektoren zukunftsweisend, wenn $v^0 > 0$ ist. Der Nullkegel zerfällt in zwei Halbkegel, den **Zukunftskegel** mit Vektoren für die $v^0 > 0$ und den **Vergangenheitskegel**, wo $v^0 < 0$ gilt.

Nullkegel

Zukunftskegel

Vergangenheitskegel

Der Nullkegel teilt den Vektorraum $\mathbb{V}$ in drei disjunkte Gebiete, das jeweilige Innere der beiden Halbkegel und deren gemeinsames Äußeres. Der erste Basisvektor selbst liegt im Inneren des Zukunftskegels, welches daher genau die zukunftsweisenden, zeitartigen Vektoren umfasst. Das Innere des Vergangenheitskegels besteht dann aus den vergangenheitsweisenden, zeitartigen Vektoren und das Außengebiet des Nullkegels schließlich besteht aus allen raumartigen Vektoren.

In der Euklidischen Geometrie spielt die Drehgruppe eine herausragende Rolle, die darin begründet ist, dass diese Gruppe die Metrik invariant lässt. In analoger Weise können wir hier nach der **Invarianzgruppe** der Metrik fragen. Wir interessieren uns also für Abbildungen $S : \mathbb{V} \rightarrow \mathbb{V}$, die (i) linear sind und (ii) die Metrik invariant lassen. Die Linearität erlaubt es uns, S als $(1, 1)$ -Tensor S^a_b zu schreiben. Wir suchen nun also Abbildungen $v^a \mapsto \hat{v}^a = S^a_b v^b$ mit

Invarianzgruppe

$$\eta_{ab}\hat{v}^a\hat{u}^b = \eta_{ab}v^au^b$$

für beliebige Vektoren u^a und v^a . Setzen wir hier die Ausdrücke für die Bildvektoren ein, ergibt sich

$$\eta_{cd} = \eta_{ab} S^a_c S^b_d$$

als Bedingung an die Abbildung S . Bezüglich einer ONB muß die Matrixdarstellung $S = (S^i_k)_{i,k=0:3}$ also die Gleichung

$$\eta_{ik} = \sum_{l,m=0}^3 \eta_{lm} S^l_i S^m_k, \quad i, k = 0 : 3$$

erfüllen. In Matrizenschreibweise lautet diese Gleichung³

$$\eta = S^t \eta S.$$

Lorentz-Gruppe Diese Relation bestimmt die Matrixdarstellung der gesuchten Abbildungen. Offensichtlich ist $(\det S)^2 = 1$, die Abbildungen also alle umkehrbar. Es wird damit folglich eine Gruppe definiert. Diese Gruppe heißt **Lorentz-Gruppe**. Sie spielt im Rahmen der Lorentz-Geometrie die gleiche Rolle wie die Drehgruppe in der euklidischen Geometrie.

Lorentz-Transformationen Die Elemente der Lorentz-Gruppe, die **Lorentz-Transformationen** haben offensichtlich die Eigenschaft, eine ONB, also ein lokales Bezugssystem, in ein anderes solches abzubilden. Darin liegt die *physikalische Bedeutung* der Lorentz-Gruppe. Sie bildet (im Rahmen der Speziellen Relativitätstheorie) Inertialsysteme auf Inertialsysteme ab.

4.2 Physikalische Interpretation

invariantes Abstandsquadrat Wie teilt sich diese Struktur im Minkowski-Vektorraum nun der Minkowski-Raumzeit selbst mit? Zu je zwei Ereignissen P und Q in $\mathbb{M}$ gibt es genau einen Vektor in $\mathbb{V}$. Das Betragsquadrat $\Delta(P, Q)$ dieses Vektors heißt **invariantes Abstandsquadrat** von P und Q . Man beachte, dass dies eine Verallgemeinerung des üblichen Abstandsbegriffs ist, da das invariante Abstandsquadrat positiv, negativ aber auch Null werden kann. Der Vektor zwischen P und Q fällt in genau eine der fünf verschiedenen Klassen. Man nennt daher die Ereignisse P und Q zeitartig, raumartig oder lichtartig getrennt, wenn ihr Verbindungsvektor zeit-, raum- oder lichtartig ist. Entsprechend sagt man 'P geht Q' voraus, bzw. 'Q folgt P', wenn der Vektor von P nach Q zukunftsgerichtet ist.

Lichtkegel Wählen wir uns ein Ereignis O als Ursprung, dann lassen sich alle anderen Ereignisse P als Ortsvektoren bzgl. O darstellen. Von besonderem Interesse sind die Ereignisse, die von O lichtartig getrennt sind. Sie bilden einen Kegel mit Vertex in O , den **Lichtkegel**⁴ von O . Die zeitartig von O getrennten Ereignisse liegen im Inneren des Lichtkegels,

³Dies ist zu vergleichen mit der definierenden Relation $S^t S = 1$ für die Drehgruppe im Euklidischen Vektorraum.

⁴Man beachte, der *Nullkegel* besteht aus Vektoren, der *Lichtkegel* aus Ereignissen.

die raumartig getrennten Ereignisse im Äußeren. Da wir jedes Ereignis als Ursprung nehmen dürfen, befindet sich somit *an jedem Ereignis* der Raumzeit ein Lichtkegel.

Wie ist dieser Kegel nun physikalisch zu interpretieren? Die gängige Vorstellung ist die eines Lichtsignals, ausgehend von einer punktförmigen Quelle, welches im Ereignis P 'gezündet' wird. Von P ausgehend breitet sich eine Wellenfront aus und überstreicht eine Menge von Raumzeitpunkten, je 'weiter weg' umso 'später', weil sich das Signal nicht unendlich schnell fortpflanzt, sondern eine endliche Ausbreitungsgeschwindigkeit besitzt. Die Ereignisse 'Empfang des Lichtsignals' bilden also einen Kegel, den wir mit dem Lichtkegel von P identifizieren. Der Lichtkegel besteht also aus genau den Ereignissen, die von P aus mit Lichtsignalen erreichbar sind. Hier ist es nicht wesentlich, dass es sich um Licht als Medium handelt. Von Bedeutung ist lediglich, dass es eine maximale Geschwindigkeit gibt, mit der sich physikalische Signale ausbreiten können. Licht hat nun gerade die Eigenschaft sich (im Vakuum) mit dieser Maximalgeschwindigkeit auszubreiten. Andere Signale (Schall, geworfene Steine,...) ausgehend von P breiten sich langsamer aus. Die jeweiligen Ereignisse Q 'Signal empfangen' folgen dem Ereignis P nach (Kausalität!). Sie sind von P zeitlich getrennt, liegen also im Inneren des Zukunftskegels von P. Sie bilden die **kausale Zukunft** des Ereignisses P. Umgekehrt, liegen alle Ereignisse, von denen aus man Signale senden kann, die in P ankommen, im Vergangenheitskegel von P. Sie bilden die **kausale Vergangenheit** von P.

kausale Zukunft

kausale
Vergangenheit

Ein sich bewegendes Teilchen überstreicht während seiner Bewegung verschiedene Raumzeitpunkte, die zusammen eine Linie γ bilden, die **Weltlinie** des Teilchens. Je zwei Ereignisse auf der Weltlinie sind zeitartig voneinander getrennt, denn sie werden durch ein Signal, nämlich das Teilchen selbst, miteinander verbunden. Ihr Verbindungsvektor ist daher zeitartig. Betrachten wir zu einem Ereignis O auf der Weltlinie eine Folge von späteren Ereignissen Q_λ mit $O = Q_0$, dann bilden deren Ortsvektoren bzgl. O eine Folge von zukunftsweisenden, zeitartigen Vektoren $\text{vec}(OQ_\lambda) = q^a(\lambda)$. Falls die Linie glatt genug ist (was wir immer annehmen), dann ist auch der Grenzübergang

Weltlinie

$$\lim_{\lambda \rightarrow 0} \frac{1}{\lambda} q^a(\lambda)$$

definiert und zukunftsweisend zeit- oder lichtartig. Eine Weltlinie γ ist daher immer eine **kausale Kurve**, in dem Sinne, dass ihr Tangentenvektor in jedem Ereignis zukunftsweisend zeit- oder lichtartig ist.

kausale Kurve

Dies ist die Interpretation, die man einer Weltlinie unabhängig von der Art und Weise, wie die Kurve durchlaufen wird, also als Menge von Ereignissen, zusprechen kann. Sie beruht auf der Existenz des Nullkegels, der durch die Metrik festgelegt ist. Der invariante Abstand definiert aber nicht nur den Kegel, sondern zusätzlich auch eine Skala. Es gibt viele Parametrisierungen für die gleiche Weltlinie, es ist jedoch eine darunter ausgezeichnet. Hier kommt noch einmal die Metrik ins Spiel.

Zwei zeitlich getrennte Ereignisse P und Q haben positives invariantes Abstandsquadrat $\Delta(P, Q)$ und wir nennen $\tau(P, Q) = \sqrt{\Delta(P, Q)}$ das invariante Zeitintervall zwischen

den beiden Ereignissen. Sind nun O und Q_λ Ereignisse auf einer Weltlinie wie oben, dann können wir zu jedem Wert des Parameters λ die Zahl $\tau(\lambda) := \tau(O, Q_\lambda)$ bilden. Wir betrachten nun die Verhältnisse $\tau(\lambda) : \lambda$, die angeben, wie sich das invariante Zeitintervall zwischen O und Q_λ zum jeweiligen Parameterwert λ verhält. Im Grenzfall bilden wir

$$\zeta = \lim_{\lambda \rightarrow 0} \frac{\tau(\lambda)}{\lambda} = \frac{d\tau}{d\lambda}(0).$$

Dieser Grenzwert bedeutet, dass der Parameter λ im Ereignis O mit dem invarianten Zeitintervall im Verhältnis ζ steht. Ist $\zeta = 1$, dann laufen die beiden Parameter τ und λ *momentan* synchron, d.h. λ misst momentan die invariante Zeitdauer zwischen O und einem infinitesimal benachbarten Ereignis auf der Weltlinie. Gilt $\zeta = 1$ auf der ganzen Weltlinie, dann nennt man den Parameter λ **Eigenzeit**. Die Vorstellung ist dabei die, dass auf der Weltlinie eine Uhr mitläuft, deren Zeitmessung als Parameter verwendet wird. Man verwendet üblicherweise die Bezeichnung τ für die Eigenzeit einer Weltlinie.

Ist eine Weltlinie mit der Eigenzeit τ parametrisiert, so folgt leicht, dass ihr Tangentenvektor u^a zu jedem Zeitpunkt *normiert* ist

$$u_a u^a = \eta_{ab} u^a u^b = 1.$$

Übung 4.3: Zeigen Sie, dass der Tangentenvektor einer mit Eigenzeit parametrisierten Weltlinie normiert ist.

Der so definierte Tangentenvektor heisst (momentane) **4-Geschwindigkeit** und dementsprechend ist die Ableitung der 4-Geschwindigkeit nach der Eigenzeit, also ihre momentane Änderung, die **4-Beschleunigung** b^a . Es gilt also

$$b^a = \dot{u}^a.$$

Die Normierungsbedingung an die 4-Geschwindigkeit hat zur Folge, dass

$$0 = \frac{d}{d\tau} (\eta_{ab} u^a u^b) = 2 \eta_{ab} u^a \dot{u}^b.$$

Die 4-Beschleunigung steht also immer senkrecht auf der 4-Geschwindigkeit, ist daher immer ein raumartiger Vektor.

Die gesamte Eigenzeitdauer $T(P, Q)$, die zwischen zwei Ereignissen P und Q entlang einer Weltlinie γ vergeht, bekommt man durch Integration entlang der Weltlinie

$$T(P, Q) = \int_\gamma d\tau.$$

Dies ist ganz in Analogie zur Definition der Bogenlänge entlang einer Kurve im euklidischen Raum. Ebenso wie dort die Bogenlänge von der Kurve zwischen zwei Punkten abhängig ist, so ist auch hier die Eigenzeit zwischen zwei Ereignissen davon abhängig,

wie die Ereignisse verbunden werden. Im euklidischen Fall ist die gerade Linie zwischen zwei Punkten dadurch ausgezeichnet, dass ihre Länge minimal ist, alle anderen Kurven zwischen den gleichen Punkten sind länger, sie machen einen 'Umweg'. Im Fall von Weltlinien in der Minkowski-Raumzeit ist die gerade Verbindungslinie zwischen zwei Ereignissen ebenfalls ausgezeichnet, jedoch dadurch dass alle 'Umzeiten', also Weltlinien zwischen den gleichen Punkten, kleiner sind als die Eigenzeit entlang der geraden Linie (das ist die Länge ihres Verbindungsvektors). Dies ist eine Form des Effekts der **Zeit-Dilatation** (vgl. unten) und stellt eine Auflösung des Zwillingsparadoxons dar. Die Tatsache folgt aus der folgenden

Zeit-Dilatation

Übung 4.4: Zeigen Sie die 'umgekehrte' Dreiecksungleichung: sind u^a und v^a zwei zukunftsgerichtete, zeitartige Vektoren dann ist auch $w^a = u^a + v^a$ zukunftsgerichtet, zeitartig, und es gilt

$$\sqrt{w_a w^a} \geq \sqrt{u_a u^a} + \sqrt{v_a v^a}.$$

Eine Theorie ist nur dann eine 'gute', physikalische Theorie, wenn sie eine Verbindung zum Experiment herstellt⁵. Sie muss darlegen, welche Größen messbar sind und erklären, unter welchen Umständen ein Experimentator welche Daten misst. Im vorliegenden Fall der Speziellen Relativitätstheorie dreht es sich um Raumzeit-Phänomene und wir müssen uns daher zunächst einmal Gedanken machen, wie ein Experimentator diese beschreiben wird. Dazu benötigen wir den Begriff eines Beobachters und seines Bezugssystems.

Wir stellen uns also einen Physiker vor, der in seinem Labor Raumzeitphänomene, also Bewegungen, beobachtet. Um die Bewegung zu beschreiben, braucht er eine Uhr und einen Maßstab. Zusätzlich wird er einen Ursprungsort und drei Richtungen festlegen, die senkrecht aufeinander sind. Dies soll dazu dienen, jeden Ort im Labor eindeutig durch die Angabe von drei Zahlen (z.B. den Abstand des Orts von der Ebene durch den Ursprung senkrecht zu den jeweiligen Richtungen, oder die Länge der Projektionen des Ortsvektors auf die jeweiligen Richtungen) zu charakterisieren⁶. Die Bewegung kann dann dadurch 'katalogisiert' werden, dass man *ab einem gewissen Zeitpunkt* zu jedem folgenden Zeitpunkt (der mit der Uhr gemessen wird) die drei Koordinaten des Orts der Bewegung registriert.

Wir sehen also hier die notwendigen 'messtechnischen' Zutaten, die zu einer Bewegungsbeschreibung notwendig sind: ein *Ursprungsereignis*, hier das Ereignis 'Uhr wird im räumlichen Ursprung gestartet', *drei räumliche Richtungen*, die wir der Einfachheit halber senkrecht zueinander wählen, *eine zeitliche Richtung*, gegeben durch den Gang der Uhr und je einen räumlichen und zeitlichen *Einheitsmaßstab*.

Im Rahmen der Speziellen (ebenso wie der Allgemeinen) Relativitätstheorie werden diese physikalischen Begriffe zum mathematischen Begriff des lokalen Bezugssystems vereint. Ein **lokales Bezugssystem** besteht aus einem Ursprungsereignis O in der

lokales
Bezugssystem

⁵was den heutigen 'theories of everything', wie z.B. der String-Theorie total abgeht.

⁶Es ist hier nicht notwendig, dass die Richtungen senkrecht aufeinander gewählt werden. Dies führt nur auf eine rechnerische Vereinfachung. Wesentlich ist, dass man weiss, wann zwei Richtungen senkrecht aufeinander sind.

Minkowski-Raumzeit und einem Quadrupel von Einheitsvektoren, einer zeitartig, die anderen raumartig, die im Sinne der Minkowski-Metrik η_{ab} senkrecht aufeinander sind. Zur Wahl der Einheiten verabreden wir, die Einheitslänge und -zeit dadurch zu verknüpfen, dass wir als Einheitslänge die Länge derjenigen Strecke zu nehmen, die das Licht innerhalb einer Zeiteinheit zurücklegt.

Beobachter Auch für einen Experimentator vergeht die Zeit. Er und sein Labor bewegen sich durch die Raumzeit. Insbesondere bilden die Raumzeitereignisse, die der räumliche Ursprung durchläuft, eine Weltlinie. An jedem Ereignis dieser Weltlinie befinden sich die drei Raumachsen des Labors und bilden zusammen mit der Zeitrichtung ein lokales Bezugssystem. Die Zeitrichtung ist durch den Gang einer Uhr gegeben, die sich im räumlichen Ursprung befindet. Zu jedem Ereignis auf der Weltlinie ist sie daher durch einen Tangentenvektor an die Weltlinie gegeben. Weil die Uhr ihre Eigenzeit anzeigt, ist der Tangentenvektor normiert. Wir kommen so zum Begriff **Beobachter**. Dies ist eine Weltlinie, durch Eigenzeit parametrisiert, an der sich zu jedem Ereignis ein lokales Bezugssystem befindet, dessen zeitliche Achse die 4-Geschwindigkeit der Weltlinie ist.

Man beachte, dass wir hier keine Einschränkungen weder an die Form der Weltlinie noch an die Orientierung der Achsen in aufeinander folgenden Ereignissen machen. Physikalisch bedeutet dies, dass wir es zulassen, dass sich das Labor beschleunigt bewegen und auch seine Orientierung ändern darf.

inertiale Beobachter Eine ausgezeichnete Klasse von Beobachtern sind daher diejenigen, die sich *nicht* beschleunigt bewegen und deren Achsen *nicht* rotieren. Dies sind **inertiale Beobachter**. Im Rahmen der SRT ist ein solcher Beobachter dadurch charakterisiert, dass die Weltlinie eine zeitartige Gerade ist und dass man die räumlichen Achsen von einem Ereignis zum nächsten durch eine zeitliche Translation entlang der 4-Geschwindigkeit u^a aufeinander abbilden kann. Dadurch sind alle lokalen Bezugssysteme entlang der Weltlinie des Beobachters festgelegt, wenn man es an einem Punkt der Weltlinie festlegt. Mithilfe von räumlichen und zeitlichen Translationen lassen sich zu einer gegebenen solchen Weltlinie beliebig viele *parallele* Weltlinien generieren (Beobachter 'weiter drüben'), so dass schließlich durch jedes Raumzeit-Ereignis eine Weltlinie läuft und insbesondere an jedem Ereignis ein lokales Bezugssystem sitzt. Wählen wir daher an einem Ereignis einmal ein lokales Bezugssystem, dann können wir an jedem anderen Ereignis ein dazu paralleles Bezugssystem festlegen. Durch diese Auswahl eines lokalen Bezugssystems an einem Ereignis wird an jedem anderen Ereignis 'das gleiche' Bezugssystem festgelegt und damit ein globales System definiert, ein **Inertialsystem**.

Die spezielle Relativitätstheorie wird vornehmlich bzgl. solcher Inertialsysteme formuliert. Der Name 'Relativitätstheorie' macht klar, dass es in dieser Theorie um die relative Beziehung zwischen Beobachtern geht. Der Zusatz 'speziell' bedeutet, dass es sich zunächst nur um eine eingeschränkte Klasse von Beobachtern handelt, die inertialen Beobachter. Diese Einschränkung wird in der allgemeinen Theorie aufgehoben. Diese Theorie ist nicht etwa subjektiv, also beobachterabhängig. Vielmehr wird innerhalb der Theorie eine objektive Struktur eingeführt, die Raumzeit, und dann diese Struktur und die physikalischen Phänomene, die sich innerhalb dieser Struktur objektiv vollziehen

unter dem subjektiven Blickwinkel verschiedener Beobachter beschrieben und interpretiert. Die objektive, beobachterunabhängige Existenz dieser Struktur hat zur Folge, dass man bestimmen kann, was ein Beobachter misst, wenn man weiß, was ein anderer Beobachter misst. Beobachterabhängig ist nur die Beschreibung und Interpretation der Phänomene, nicht jedoch das Phänomen an sich.

Bemerkung Wie ist das Prinzip von der Konstanz der Lichtgeschwindigkeit in diesem geometrischen Rahmen verwirklicht? Die Unabhängigkeit der Ausbreitungsgeschwindigkeit des Lichts von der Quelle bedeutet, dass der Lichtkegel, der vom Ereignis der Lichtemission ausgeht, allein durch das Licht und sein Medium, in dem es sich ausbreitet (hier das Vakuum, also die reine Raumzeitstruktur) bestimmt ist. Das bedeutet insbesondere, dass der Lichtkegel eines Ereignisses nur durch das Ereignis selbst bestimmt ist und nicht etwa durch einen Beobachter (der auf der Lichtquelle sitzt) beeinflusst wird. Wäre dies so, dann gäbe es am gleichen Ereignis u.U. für jeden Beobachter einen anderen Lichtkegel. Der Lichtkegel wäre abhängig von der 4-Geschwindigkeit der Beobachter. Dies führt dann nicht zur Lorentz-Geometrie, sondern zu einer sog. **Finsler-Geometrie**. Die sich daraus ergebenden Effekte sind (bisher) nicht beobachtet.

Finsler-Geometrie

Wie haben wir uns die Raum- und Zeitmessung eines Beobachters vorzustellen? Dies ist in idealisierter Form recht einfach und geht schon aus der Beschreibung des Bezugssystems hervor. Nehmen wir an, wir wollen ein Ereignis P bzgl. des Ursprungs O unseres Bezugssystems charakterisieren, dann müssen wir angeben, 'wo' und 'wann' es stattfindet. Wir müssen also separate Raum- und Zeitmessungen durchführen. *Ein Beobachter zerlegt die Raumzeit in Raum und Zeit*. Mathematisch bedeutet dies, dass der Ortsvektor p^a des Ereignisses P in einen räumlichen und einen zeitlichen Anteil zerlegt wird

$$p^a = \Delta\tau u^a + x^a,$$

dabei ist $u_a x^a = 0$, also x^a raumartig. Auf diese Weise wird jeder Ortsvektor *eindeutig* in einen Anteil parallel zur 4-Geschwindigkeit u^a und einen Anteil senkrecht dazu zerlegt. Die Größe $\Delta l = \sqrt{-\eta_{ab} x^a x^b}$ ist die Länge des räumlichen Anteils, also die räumliche Distanz, die der Beobachter den beiden Ereignissen P und O zumisst. Und analog dazu ist die 'Länge' $\Delta\tau$ des zeitlichen Anteils gerade die Eigenzeit, die für den Beobachter zwischen O und P verstreicht. Es gilt offensichtlich

$$\Delta\tau = u_a p^a.$$

Für Vektoren mit $\Delta\tau = 0$, also solche, die senkrecht auf u^a stehen, gibt es keine Eigenzeitdifferenz, die zugehörigen Ereignisse sind demnach *gleichzeitig* mit O . Man nennt daher den Unterraum von $\mathbb{V}$ bzw. die Hyperebene durch O senkrecht zu u^a auch den **Raum der Gleichzeitigkeit** bzgl. u^a .

Raum der
Gleichzeitigkeit

Verschiedene Beobachter haben verschiedene 4-Geschwindigkeiten. Daher ist die Zerlegung der Vektoren in Raum- und Zeitanteil verschieden. Insbesondere sind die Räume der Gleichzeitigkeit verschieden. Das bedeutet, dass verschiedene Beobachter eine

verschiedene Wahrnehmung haben hinsichtlich der Gleichzeitigkeit zweier Ereignisse.

Wir betrachten nun die Welt durch die 'Augen' zweier inertialer Beobachter, 'Alice' und 'Bob' (vgl. Abb. 4.1). Wir nehmen an, dass die beiden das gleiche Ursprungsereignis O wählen. Die beiden Weltlinien sind Geraden durch O mit jeweiligen Tangentenvektoren

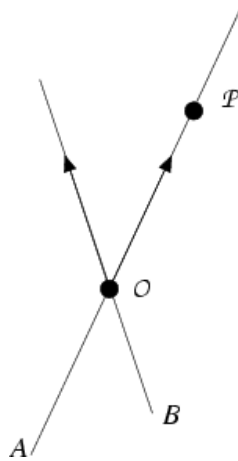


Abbildung 4.1: Zur Zeit-Dilatation

u_A^a und u_B^a . Alice trage mit sich eine Uhr, die sie nach einer Eigenzeit von $\Delta\tau_A$ abliest. Dies definiert ein Ereignis P in der Raumzeit $\mathbb{M}$ und wir fragen nun, wie dieser Vorgang von Bob beobachtet und interpretiert wird. Da sich Alice bzgl. Bob bewegt, ist auch die Uhr, die sie trägt bewegt. Um die Bewegung der Uhr beschreiben zu können, brauchen wir den Ortsvektor p^a von P . Offensichtlich ist

$$p^a = \Delta\tau_A u_A^a.$$

Bob spaltet nun diesen Vektor in Raum- und Zeitanteile auf und erhält

$$p^a = \Delta\tau_B u_B^a + x^a,$$

wobei $x^a \neq 0$ ist. Für Bob ist die Uhr in gleichförmiger Bewegung. Sie hat zwischen O und P nach seiner Messung innerhalb der Zeit $\Delta\tau_B$ die Strecke x^a zurückgelegt. Er ermittelt daher ihre Geschwindigkeit, die **Relativgeschwindigkeit** zwischen ihr und ihm, zu

$$v^a = \frac{x^a}{\Delta\tau_B}.$$

Das bedeutet insbesondere, dass für Bob die zwischen O und P verstrichene Zeit eine andere ist als für Alice

$$\Delta\tau_B = \eta_{ab} p^a u_B^b = \gamma \Delta\tau_A,$$

Lorentz-Faktor wobei $\gamma = \eta_{ab} u_A^a u_B^b$ der **Lorentz-Faktor** zwischen den beiden 4-Geschwindigkeiten

ist. Da es sich um 4-Geschwindigkeiten handelt, also zeitartige, zukunftsgerichtete Einheitsvektoren folgt aufgrund der umgekehrten Dreiecksungleichung $\gamma \geq 1$. Wir erhalten daher

$$\Delta\tau_B \geq \Delta\tau_A$$

und damit den bekannten Effekt der **Zeit-Dilatation**: *Die bewegte Uhr geht langsamer* (denn sie zeigt eine kürzere Zeit an). Man beachte, dass es sich hierbei um einen vollkommen symmetrischen Effekt handelt. Man überzeugt sich leicht, dass auch Alice eine von Bob mitgenommene Uhr als langsamer als ihre eigene beurteilt.

Zeit-Dilatation

Kehren wir zu Bobs Raum-Zeit-Zerlegung des Ortsvektors p^a zurück

$$p^a = \Delta\tau_A u_A^a = \Delta\tau_B u_B^a + x^a = \Delta\tau_B (u_B^a + v^a),$$

woraus sich für Alices 4-Geschwindigkeit der Ausdruck

$$u_A^a = \gamma (u_B^a + v^a)$$

ergibt. Aus der Normierungsbedingung folgt $1 = \gamma^2(1 + v_a v^a)$. Da v^a ein raumartiger Vektor ist, gilt $0 \geq v^a v_a = -v^2$, so dass sich der Lorentz-Faktor aus der Geschwindigkeit der bewegten Uhr nach der bekannten Formel

$$\gamma = \frac{1}{\sqrt{1 - v^2}}$$

berechnen lässt.

Wir betrachten einen ausgedehnten Körper, der sich gleichförmig durch die Raumzeit bewegt. Wir nehmen der Einfachheit halber an, dass der Körper nur in einer Dimension eine nennenswerte Ausdehnung besitzt und in den anderen beiden Dimensionen vernachlässigbar 'dünn' sei. Jeder Punkt auf diesem 'Stab', insbesondere Anfangs- und Endpunkte, beschreibt eine Gerade in der Raumzeit. Nehmen wir ferner an, Alice bewege sich mit dem Stab, der dann also in ihrem Bezugssystem ruht (vgl. Abb. 4.2). Will Alice die Ausdehnung des Stabes messen, muss sie die räumliche Distanz zwischen Anfangs- und Endpunkt des Stabes bestimmen. Sie nimmt also zwei Ereignisse, je eines auf der jeweiligen Weltlinie, die für sie gleichzeitig stattfinden, also z.B. O und P. Der Vektor p^a zwischen O und P ist raumartig und seine Länge ist die Länge des Stabes, wenn Alice ihn misst:

$$\Delta l_A = \sqrt{-\eta_{ab} p^a p^b}.$$

Wenn Bob die Länge des Stabes misst, geht er genauso vor: er bestimmt das Ereignis Q auf der Endlinie des Stabes, welches mit O *für ihn* gleichzeitig ist und berechnet die Länge des Ortsvektors q^a .

Mathematisch betrachtet müssen wir, um die von einem Beobachter gemessene Länge des Stabes zu bestimmen, folgendermaßen vorgehen. Wir wählen O als Ursprung und haben nun ein Ereignis auf der Endlinie des Stabes zu bestimmen, das für den Beobachter gleichzeitig mit O stattfindet. Das heißt, wir müssen einen Ortsvektor bestimmen,

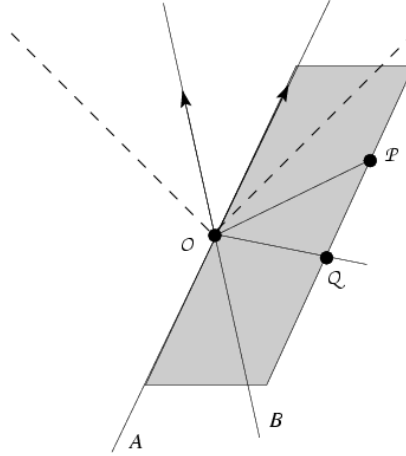


Abbildung 4.2: Zur Lorentz-Kontraktion

der senkrecht auf der 4-Geschwindigkeit u^a des Beobachters ist und auf ein Ereignis auf der Endlinie zeigt. Im vorliegenden Fall ist die Endlinie eine Gerade, deren Tangentenvektor die 4-Geschwindigkeit u_A^a von Alice ist, da der Stab in ihrem System ruht. Sie besteht aus allen Ereignissen, deren Ortsvektor x^a die Form

$$x^a = s^a + \lambda u_A^a$$

für beliebige reelle Zahlen λ hat. Dabei ist s^a der Ortsvektor zu einem beliebigen Ereignis auf der Linie. Wir können also insbesondere hier den Ortsvektor zu P einsetzen.

Wollen wir nun die von Bob gemessene Länge des Stabs bestimmen, benötigen wir den Ortsvektor q^a zum Ereignis Q auf der Endlinie, das mit O gleichzeitig ist. Dieser erfüllt die Gleichung

$$0 = \eta_{ab} q^a u_B^b.$$

Einsetzen der Parameterdarstellung für die Endlinie ergibt

$$q^a = p^a - \frac{p_c u_B^c}{\gamma} u_A^a$$

und damit die Relation zwischen den gemessenen Längen:

$$(\Delta l_B)^2 = (\Delta l_A)^2 - \left(\frac{p_c u_B^c}{\gamma} \right)^2.$$

Es gilt also immer $\Delta l_B \leq \Delta l_A$, d.h. *der bewegte Maßstab ist kürzer*. Dies ist der bekannte Effekt der **Lorentz-Kontraktion**.

Lorentz-Kontraktion

In dem speziellen Fall, in dem sich die Beobachter 'entlang des Stabes' bewegen, wenn also der Stab in der zwei-dimensionalen Ebene liegt, die von den beiden 4-Geschwindigkeiten aufgespannt wird, lässt sich diese Formel noch etwas vereinfachen.

Übung 4.5: Zeigen Sie, dass sich die bekanntere Formel

$$\Delta l_B = \frac{1}{\gamma} \Delta l_A$$

ergibt.

Man beachte: die Zeit-Dilatation und die Lorentz-Kontraktion sind beobachterabhängige Effekte. Sie kommen dadurch zustande, dass verschiedene Beobachter unterschiedliche 'Blickwinkel' auf das gleiche objektive Raumzeit-Phänomen (z.B. Ablesen einer Uhr, Stab in der Raumzeit) haben.

Auch in der Beschreibung von Bewegung gibt es beobachterabhängige Unterschiede. Insbesondere ist die Geschwindigkeit mit der sich eine Bewegung vollzieht beobachterabhängig. Die geometrische Formulierung der RT gestattet auch hier eine einfache geometrische Analyse der Zusammenhänge, die hier aber zu weit führen würde. Am Ende dieser Analyse steht die wohlbekannte Formel für die 'Addition von Geschwindigkeiten', das **Additionstheorem** für Geschwindigkeiten.

Additionstheorem

Aus den Betrachtungen in diesem Kapitel wollen wir folgende Einsichten mitnehmen:

1. Die Menge aller Ereignisse lässt sich mit hoher Genauigkeit als ein 4-dimensionales Kontinuum beschreiben. Jeder Raumzeitpunkt, jedes Ereignis, wird durch die Angabe von vier reellen Zahlen charakterisiert.
2. Dieser Ereignismenge werden durch eine Lorentz-Metrik Raumeigenschaft aufgedrückt: an jedem Ereignis gibt es einen Lichtkegel, und eine Einheitszeit (und damit -länge). Die Metrik definiert einen invarianten Abstandsbegriff.
3. In einem lokalen Inertialsystem hat die Matrixdarstellung der Metrik die Standard-Form

$$\begin{pmatrix} +1 & & & \\ & -1 & & \\ & & -1 & \\ & & & -1 \end{pmatrix}$$

4. Ein kräftefreier Körper bewegt sich auf einer geraden Weltlinie.

Diese Eigenschaften haben alle *lokalen Charakter* in dem Sinne, dass sie sich nur an die Struktur an einem einzigen (aber beliebigen) Ereignis wenden. Dies ist für die letzte Eigenschaft nicht ganz offensichtlich, aber doch nachvollziehbar. Denn wie kann man sich eine gerade Linie vorstellen? Doch so, dass es eine Linie ist, die von ihrer Richtung nicht abweicht. D.h., aus einer Richtung an einem Ereignis angekommen, setzt sie ihren Weg *in der gleichen Richtung* fort. Die Änderung ihres Tangentenvektors an diesem Ereignis ist demnach proportional zum Tangentenvektor selber. Eine gerade Linie ist also auch durch ihre Eigenschaften an einem einzigen Ereignis zu charakterisieren.

Die Beschränkung auf lokale Eigenschaften vermeidet den Begriff eines globalen Inertialsystems. Wir hatten schon gesehen, dass dessen Definition fragwürdig ist und die

Existenz eines globalen Inertialsystems letztlich postuliert werden muss. Im Rahmen einer Gravitationstheorie kann dieses Postulat nicht mehr aufrecht erhalten werden. Damit fällt auch die Minkowski-Raumzeit als allgemein gültiges Modell für die Raumzeit weg, denn erst das Postulat der Existenz eines Inertialsystems zusammen mit dem Relativitätsprinzip führte uns auf dieses Modell eines affinen Raums.

5 Die gekrümmte Raumzeit

5.1 Die Raumzeit als Mannigfaltigkeit

Wir hatten gesehen, dass die Raumzeit der SRT ein affiner Raum ist. Das heißt, sie besteht aus Ereignissen mit jeweils daran angeheftetem Vektorraum von Ortsvektoren. Je zwei dieser Vektorräume, die sich an verschiedenen Ereignissen befinden, lassen sich durch eine Parallelverschiebung miteinander identifizieren. Wesentlich bei dieser Identifikation ist dabei die folgende Tatsache. Wir betrachten vier Ereignisse O , P , P' und O' mit den jeweiligen Vektorräumen (vgl. Abb. 5.1) Nun verschieben wir einen Vektor

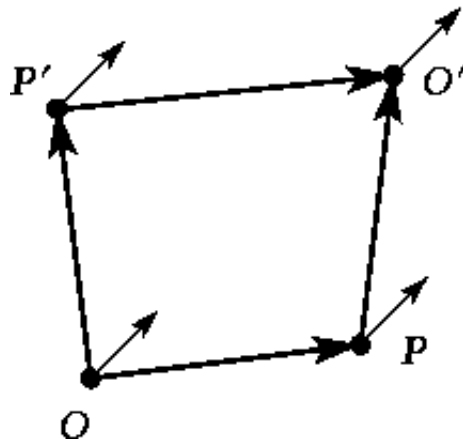


Abbildung 5.1: Parallelverschiebung ist wegunabhängig

im Vektorraum bei O zuerst nach P und dann nach O' . Anschließend verschieben wir denselben Vektor bei O zuerst nach P' und dann nach O' . Wir erhalten das gleiche Resultat. Das bedeutet, dass es bei der Parallelverschiebung von O nach O' nicht auf die Zwischenstationen ankommt, sondern nur auf Anfangs- und Endereignis.

Übung 5.1: Zeigen Sie, dass diese Eigenschaft aus den Grundeigenschaften für affine Räume folgt.

In diesem Sinne sind alle Vektorräume der speziell-relativistischen Raumzeit 'parallel' zueinander. Die Raumzeit ist flach.

Auf einer gekrümmten Fläche ist Parallelverschiebung eines Vektors nicht mehr notwendigerweise unabhängig vom Weg. Wir betrachten dazu als Beispiel eine Kugeloberfläche (vgl. Abb. 5.2). Offensichtlich ist die Parallelverschiebung eines Vektors auf der

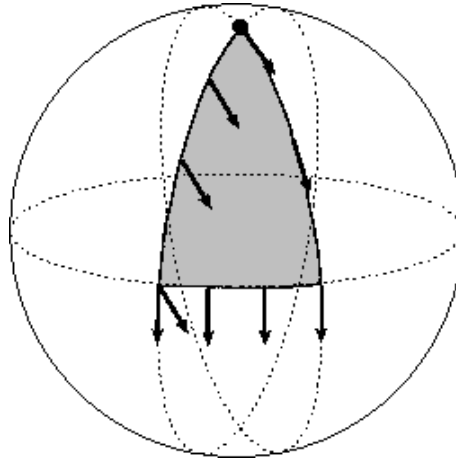


Abbildung 5.2: Parallelverschiebung auf der Kugel

Kugel ausgehend vom Nordpol zu einem Punkt auf dem Äquator nicht mehr unabhängig davon, wie man die Punkte verbindet. Vielmehr scheint der Unterschied zwischen den Vektoren, die auf verschiedenen Wegen erreicht werden umso größer zu sein, je größer die Fläche des Dreiecks ist, das durch die Wege umfasst wird.

5.1.1 Mannigfaltigkeiten

Bevor wir jedoch zu den Konsequenzen aus dieser Nichteindeutigkeit des Paralleltransports kommen, müssen wir uns zuerst klar werden, wie wir uns den Vektorraum vorzustellen haben, der an einem Punkt einer gekrümmten Fläche sitzt. Die intuitive Idee dabei ist diejenige, dass dieser Vektorraum einer Tangentialebene an eine Fläche im euklidischen Raum entsprechen sollte, also eine 'Approximation erster Ordnung' der Fläche sein sollte (vgl. Abb 5.3) Jedoch kann die Situation, die in dieser Abbildung

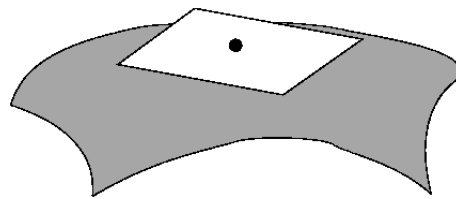


Abbildung 5.3: Tangentialebene an eine Fläche

dargestellt ist, nicht zufrieden stellen, denn offensichtlich erstreckt sich die Tangentialebene 'außerhalb' der Fläche in den umgebenden Raum. Wollen wir jedoch z.B. unser Universum als eine gekrümmte Fläche beschreiben, so fragt sich, wohinein sich die Tangentialebene erstrecken wollte? Dieses Bild hat also direkt übernommen keinen Sinn. Vielmehr müssen wir versuchen, mithilfe von intrinsischen Mitteln die Eigenschaften

von Tangentialvektoren zu beschreiben. Es geht also darum, Eigenschaften der Ortsvektoren bzw. von Tangentialvektoren an eine Fläche im euklidischen Raum zu finden, die allein durch intrinsische Eigenschaften charakterisiert werden können.

Betrachten wir also ein konkretes Beispiel, eine Fläche M im dreidimensionalen euklidischen Raum z.B. eine Kugeloberfläche. Zunächst brauchen wir eine Möglichkeit, verschiedene Punkte auf der Fläche zu unterscheiden. Dies geschieht durch eine Zuordnung von je zwei Zahlen $(x^1[P], x^2[P])$ zu einem Punkt P . Ist diese Zuordnung so, dass je zwei verschiedenen Punkten verschiedene Zahlenpaare zugeordnet werden, so spricht man von einem **Koordinatensystem** bzw. einer **Karte** für die Fläche M . Eine Karte ist also eine Abbildung von M in den $\mathbb{R}^2$. Das Beispiel der Kugel zeigt uns, dass es im allgemeinen nicht möglich ist, die gesamte Kugel durch eine einzige Karte zu überdecken. Vielmehr braucht man einen **Atlas**, eine Ansammlung von Karten, um die gesamte Kugel zu erfassen.

Koordinatensystem

Karte

Atlas

Nun muss natürlich gewährleistet sein, dass man Punkte, die in verschiedenen Karten vorkommen, eindeutig wiederfinden kann. Das heißt, jedem Punkt P in einem solchen **Überlappgebiet** zweier Karten kann man zwei Zahlenpaare $(x^1[P], x^2[P])$ und $(y^1[P], y^2[P])$ zuordnen. Damit erhält man eine Abbildung zwischen den Koordinatenpaaren von Punkten im Überlappgebiet

Überlappgebiet

$$\phi : (x^1, x^2) \mapsto P \mapsto (y^1, y^2).$$

Diese Übergangsabbildung zwischen zwei Karten ist eine Abbildung zwischen Teilmengen des $\mathbb{R}^2$, eine so genannte **Koordinatentransformation**. Für eine Kugeloberfläche findet man immer Karten derart, dass die Übergangsabbildung *differenzierbar* ist. Man findet also die Glattheitseigenschaften einer Fläche in den Koordinatentransformationen zwischen verschiedenen Koordinatensystemen wieder.

Koordinatentransformation

Bei der Definition einer n -dimensionalen differenzierbaren Mannigfaltigkeit dreht man den Spieß um. Ein topologischer Raum M ist eine n -dimensionale Mannigfaltigkeit, wenn jeder seiner Punkte in einer n -dimensionalen Karte liegt, d.h., wenn M vollständig mit Karten überdeckt werden kann, so dass jedem Punkt mindestens ein n -Tupel von Koordinaten zugeordnet werden kann. Zudem fordert man, dass die Übergangsabbildungen zwischen je zwei Karten differenzierbar sein sollen. Man kann also zu jeder Koordinatentransformation

$$y^i = \phi^i(x^k), \quad y = \phi(x)$$

die zugehörige Jacobi-Matrix

$$D\phi(x) = \left(\frac{\partial y^i}{\partial x^k} \right)_{i,k=1:n} =: \left(\phi^i_k \right)_{i,k=1:n}$$

berechnen. Es ist üblich, anzunehmen, dass diese Abbildungen beliebig oft differenzierbar sind.

5.1.2 Der Tangentialraum

Wir betrachten wieder eine Fläche M im dreidimensionalen Raum (Abb. 5.4) und wäh-

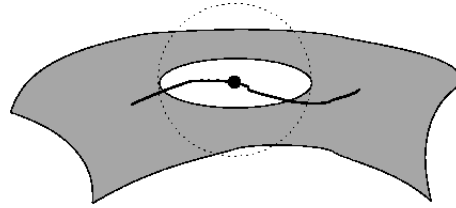


Abbildung 5.4: Umgebung eines Punktes auf M

len eine dreidimensionale Umgebung U eines Punktes von M . Nehmen wir an, in dieser Umgebung U sei eine reell-wertige Funktion f gegeben. Diese ordnet jedem Punkt in U also insbesondere jedem Punkt in $U \cap M$ einen Zahlenwert zu. Wir interessieren uns jetzt für die Änderung von f entlang einer Kurve $\gamma(\lambda)$ in M durch den Punkt $P = \gamma(0)$ auf M .

Wir bezeichnen mit (x, y, z) die kartesischen Koordinaten im dreidimensionalen Raum. Dann können wir die Kurve in der Parameterform

$$x = x(\lambda), \quad y = y(\lambda), \quad z = z(\lambda)$$

angeben, und die Werte von f auf dieser Kurve sind durch $f(\gamma(\lambda)) = f(x(\lambda), y(\lambda), z(\lambda))$ gegeben. Wir berechnen die *momentane Änderung* von f entlang γ bei P :

$$\begin{aligned} \frac{d}{d\lambda} f(\gamma(\lambda)) \Big|_{\lambda=0} &= \frac{\partial f}{\partial x}(P) \dot{x}(0) + \frac{\partial f}{\partial y}(P) \dot{y}(0) + \frac{\partial f}{\partial z}(P) \dot{z}(0) \\ &= \left(\frac{\partial f}{\partial x}(P), \frac{\partial f}{\partial y}(P), \frac{\partial f}{\partial z}(P) \right) \underbrace{\begin{pmatrix} \dot{x}(0) \\ \dot{y}(0) \\ \dot{z}(0) \end{pmatrix}}_{\mathbf{X}}. \end{aligned} \quad (5.1.1)$$

Wir stellen fest, dass die Änderung linear vom Tangentialvektor $\mathbf{X}$ an die Kurve bei P abhängt. Schneiden wir die Fläche entlang der Kurve auf, dann erhalten wir die Abb. 5.5, welche diese Tatsache noch einmal grafisch ausdrückt: die Änderung der

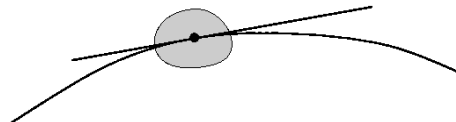


Abbildung 5.5: Umgebung eines Punktes auf M

Funktionswerte von f entlang der Kurve sind *in erster Ordnung* d.h. bis auf Terme der Ordnung λ^2 gleich der Änderung der Funktionswerte entlang der Geraden durch P in

Richtung des Tangentialvektors X . Sie hängt daher nicht von der gesamten Kurve ab, sondern nur von deren Tangentialvektor am Punkt P .

Andererseits handelt es sich bei der momentanen Änderung von f entlang γ bei P um eine *intrinsische Eigenschaft*, die ohne Bezugnahme auf den umgebenden dreidimensionalen Raum berechnet werden kann. Um dies zu sehen, nehmen wir an, dass eine Umgebung von P durch Koordinaten (x^1, x^2) überdeckt werden kann. Dann können wir die Kurve in der Form

$$x^1 = x^1(\lambda), \quad x^2 = x^2(\lambda)$$

darstellen und die Funktion f besitzt eine Koordinatendarstellung in der Form $f(x^1, x^2)$. Für die momentane Änderung von f entlang γ bei P ergibt sich nun ganz analog

$$\begin{aligned} \frac{d}{d\lambda} f(\gamma(\lambda)) \Big|_{\lambda=0} &= \frac{\partial f}{\partial x^1}(P) \dot{x}^1(0) + \frac{\partial f}{\partial x^2}(P) \dot{x}^2(0) \\ &= \left(\frac{\partial f}{\partial x^1}(P), \frac{\partial f}{\partial x^2}(P) \right) \underbrace{\begin{pmatrix} \dot{x}^1(0) \\ \dot{x}^2(0) \end{pmatrix}}_X. \end{aligned} \quad (5.1.2)$$

Offensichtlich stehen in dieser Formel keine Größen mehr, die sich auf einen umgebenden Raum beziehen; sie ist vollkommen intrinsisch.

Was sagen uns jetzt die Formeln (5.1.1) und (5.1.2)? Zunächst sehen wir wieder, dass die momentane Änderung einer Funktion f entlang einer Kurve nur von deren Tangentialvektor abhängt. Denn wenn wir eine zweite Kurve durch P betrachten, deren Tangentialvektor bei P mit X übereinstimmt, dann erhalten wir die gleiche Änderung. Bei gegebener Funktion f liefert uns ein Tangentialvektor bei P mittels dieser Formeln immer die momentane Änderung der Funktion bei P und diese hängt in linearer Weise vom Tangentialvektor ab. Ein Tangentialvektor X vermittelt eine Abbildung, die jeder Funktion f eine Zahl zuordnet

$$\hat{X} : f \mapsto \hat{X}(f) = \sum_{i=1}^2 x^i \frac{\partial f}{\partial x^i}(P),$$

wobei $X^i = \dot{x}^i(0)$. Diese Abbildung hat folgende Eigenschaften:

- (i) Sie ist *linear* in der Funktion f , d.h. für beliebige reelle Zahlen α, β und Funktionen f, g ist

$$\hat{X}(\alpha f + \beta g) = \alpha \hat{X}(f) + \beta \hat{X}(g)$$

- (ii) Sie erfüllt die *Produktregel*, d.h. für zwei Funktionen f und g gilt

$$\hat{X}(fg) = \hat{X}(f)g + f\hat{X}(g)$$

Übung 5.2: Zeigen Sie, dass aus diesen Eigenschaften $\hat{X}(c) = 0$ folgt für beliebige konstante Funktionen c .

Eine solche Abbildung nennt man **Derivation** von Funktionen. Wir finden also in unse- Derivation

rem konkreten Fall, dass jeder Tangentialvektor bei P eine Derivation von Funktionen, die in der Nähe von P definiert sind, liefert. Solche Derivationen kann man (notwendigerweise) linear kombinieren, indem man $(\alpha\hat{X} + \beta\hat{Y})f = \alpha\hat{X}(f) + \beta\hat{Y}(f)$ setzt. Dies entspricht genau der Änderung von f entlang dem Tangentialvektor $\alpha X + \beta Y$. Die Derivationen bilden also einen Vektorraum.

Soweit zum konkreten Fall. Wollen wir nun im allgemeinen Fall Tangentialvektoren an eine n -dimensionale Mannigfaltigkeit einführen, so kann dies ganz einfach geschehen. Nun können wir zwar nicht mehr sagen, ein Tangentialvektor vermittele eine Derivation, denn wir haben den Tangentialvektor im Sinne eines Vektors in einem umgebenden Raum nicht mehr zur Verfügung. Jedoch können wir jetzt einfach definieren:

Tangentialvektor **Definition 5.1.** Ein **Tangentialvektor** an M im Punkt P ist eine Derivation von Funktionen, die in der Nähe von P definiert sind.

Tangentialraum Die so definierten Tangentialvektoren bilden einen Vektorraum, den **Tangentialraum** $T_P M$ an M im Punkt P .

Wie groß ist die Dimension dieses Tangentialraums? Dazu müssen wir uns eine Basis verschaffen. Es sei $(x^i)_{i=1:n}$ ein Koordinatensystem um P . Es liegt nahe, dass die Abbildungen

$$\left. \frac{\partial}{\partial x^i} \right|_P : f \mapsto \frac{\partial f}{\partial x^i}(P)$$

Derivationen sind. Wir wollen zeigen, dass diese partiellen Ableitungen nach den Koordinaten eine Basis des Tangentialraums bei P sind. Es ist zu zeigen, dass sich jede Derivation X als Linearkombination

$$X = \sum_{i=1}^n a^i \left. \frac{\partial}{\partial x^i} \right|_P$$

darstellen lässt. Dazu nehmen wir der Einfachheit halber an, dass P die Koordinaten $x^1 = x^2 = \dots = x^n = 0$ besitzt. Außerdem beachten wir, dass jede Funktion in der Nähe von P die Darstellung

$$f(x^1, \dots, x^n) = f(0) + \sum_{i=1}^n x^i g_i(x^1, \dots, x^n)$$

besitzt.

Übung 5.3: Man zeige, wie diese Darstellung aus der Anwendung des Taylorschen Approximationssatzes folgt.

Dabei sind die g_i Funktionen in der Nähe von P , für die

$$\left. \frac{\partial f}{\partial x^i} \right|_P = g_i(0)$$

gilt. Somit ergibt sich mit den Rechenregeln (i) und (ii)

$$X(f) = \sum_{i=1}^n X(x^i) g_i(0) = \sum_{i=1}^n X(x^i) \frac{\partial f}{\partial x^i}(P)$$

also die gewünschte Darstellung für f (denn $X(x^i)$ ist eine reelle Zahl). Nun bleibt noch nachzuprüfen, dass die Koordinatenableitungen linear unabhängig sind.

Übung 5.4: Zeigen Sie dies, indem Sie eine beliebige Linearkombination sukzessive auf die Koordinatenfunktionen x^i anwenden.

Damit ergibt sich also das folgende Ergebnis

Lemma 5.1. *Der Tangentialraum an eine n -dimensionale Mannigfaltigkeit M in einem Punkt P hat die Dimension n .*

Die Basis der partiellen Ableitungsoperatoren nennt man **Koordinatenbasis** oder auch **natürliche Basis**.

Koordinatenbasis
natürliche Basis

Damit sind wir nun also so weit wie in der SRT: an jedem Punkt der Mannigfaltigkeit sitzt ein n -dimensionaler Vektorraum. Im Gegensatz zur Raumzeit-Mannigfaltigkeit der SRT ist es hier nicht offensichtlich, wie man zwei solche Vektorräume miteinander in Verbindung bringen sollte. Die dazu geeignete Struktur wird in Kürze vorgestellt. Zuerst müssen wir uns noch über die Tensoren Gedanken machen, die auf dem Tangentialraum definiert sind.

5.2 Tensorfelder

An jedem Punkt einer Mannigfaltigkeit M sitzt ein n -dimensionaler Vektorraum $\mathbb{V} = T_P M$. Ebenso wie im Kapitel 3 lassen sich nun Tensoren definieren. Der erste Schritt besteht in der Einführung des zu $\mathbb{V}$ dualen Raums $\mathbb{V}^*$. Dieser besteht aus linearen Abbildungen, die jedem Tangentialvektor in $T_P M$ in linearer Weise eine Zahl zuordnen. Wir kennen solche Abbildungen: mit jeder Funktion f ist die Abbildung, die jedem Tangentialvektor X die momentane Änderung von f in dessen Richtung, also die **Richtungsableitung** von f in Richtung X zuordnet, also

Richtungsableitung

$$df_P : X \mapsto X(f) = \sum_{i=1}^n X^i(P) \frac{\partial f}{\partial x^i}(P),$$

linear in X . Es handelt sich also um einen Kovektor. Man schreibt häufig

$$X(f) = \langle df_P, X \rangle = \sum_{i=1}^n X^i(P) \frac{\partial f}{\partial x^i}(P)$$

und nennt df das **totale Differential** von f (am Punkt P).

totale Differential

Ebenso wie die Koordinaten eine Basis im Tangentialraum definieren, so vermitteln sie auch eine Basis im **Kotangentialraum**, dem Dualraum T_p^*M von T_pM . Da jede Koordinate x^i selbst als Abbildung betrachtet werden kann, die jedem Punkt P seine i -te Koordinate zuordnet, können wir die **Koordinatendifferentiale** dx^i bilden. Diese ergeben gerade n Stück.

Übung 5.5: Man zeige, dass

$$\left\langle dx_p^i, \frac{\partial}{\partial x^k} \Big|_p \right\rangle = \delta_k^i$$

gilt, d.h., die Koordinatendifferentiale bilden die zur Basis der Koordinatenableitungen duale Basis.

Damit kann man nun das totale Differential jeder Funktion f , die in der Nähe von P definiert ist, in der Form

$$df_P = \sum_{i=1}^n a_i dx_p^i$$

schreiben.

Übung 5.6: Zeigen Sie, dass die Koeffizienten gerade durch die partiellen Ableitungen von f nach den Koordinaten gegeben sind:

$$a_i = \frac{\partial f}{\partial x^i}(P).$$

Nehmen wir nun an der betrachtete Punkt liege in einem Überlappgebiet zweier Karten mit Koordinaten $(x^i)_{i=1:n}$ bzw. $(y^k)_{k=1:n}$ mit der Koordinatentransformation $y = \phi(x)$. Dann gibt es im Tangentialraum T_pM zwei ausgezeichnete Basen, die Koordinatenableitungen $(\partial/\partial x^i|_P)_{i=1:n}$ und die $(\partial/\partial y^k|_P)_{k=1:n}$. Wie lautet die Basistransformation zwischen diesen Basen? Dazu betrachten wir eine beliebige Funktion f in der Nähe von P . Wir können diese Funktion sowohl bzgl. der x -Koordinaten als auch bzgl. der y -Koordinaten darstellen und erhalten damit notwendigerweise die Gleichung

$$f(x) = f(\phi(x))$$

in einer Umgebung von P , welches die Koordinaten x_0 bzw. y_0 besitzen möge. Daraus folgt die Gleichung

$$\frac{\partial f}{\partial x^i}(x_0) = \sum_{k=1}^n \frac{\partial f}{\partial y^k}(y_0) \frac{\partial \phi^k}{\partial x^i}(x_0),$$

so dass wir die Basistransformation in der Form

$$\frac{\partial}{\partial x^i} \Big|_P = \sum_{k=1}^n \frac{\partial y^k}{\partial x^i}(x_0) \frac{\partial}{\partial y^k} \Big|_P = \sum_{k=1}^n \phi^k_{,i}(x_0) \frac{\partial}{\partial y^k} \Big|_P. \quad (5.2.1)$$

Das heißt, die Matrix, die die Basistransformation vermittelt, ist durch die Jacobi-Matrix der Koordinatentransformation gegeben.

Die Basistransformation für die jeweiligen dualen Basen geschieht nun zwangsläufig mit der Inversen der Jacobi-Matrix (vgl. Gl. (3.1.3)). Es ist aber auch hier einfach zu sehen, wie sich dieses Transformationsverhalten ergibt. Wir betrachten die Koordinatendifferentiale $dy_p^k = d\phi_p^k$ und erhalten für diese Differentiale bzgl. der Basis der dx_p^i

$$dy_p^k = \sum_{i=1}^n \frac{\partial \phi^k}{\partial x^i}(x_0) dx_p^i = \sum_{i=1}^n \phi^k_{,i}(x_0) dx_p^i,$$

also genau das Transformationsverhalten einer dualen Basis.

Wir betrachten wir einen Tangentialvektor X am Punkt P , den wir bzgl. der beiden Basen darstellen können

$$X = \sum_{i=1}^n X^i \frac{\partial}{\partial x^i} \Big|_P = \sum_{k=1}^n \hat{X}^k \frac{\partial}{\partial y^k} \Big|_P.$$

Aus dem Transformationsverhalten der Basisvektoren erhalten wir nun das Transformationsverhalten der Komponenten von X

$$\hat{X}^k = \sum_{i=1}^n \phi^k_{,i}(x_0) X^i.$$

Dies ist genau das kontravariante Transformationsverhalten, das wir für Vektoren aus $\mathbb{V}$ in Kap. 3 kennengelernt hatten.

Übung 5.7: Zeigen Sie, dass sich die Komponenten eines totalen Differentials df in kovarianter Weise bei einer Koordinatentransformation ändern.

Damit sind unsere Tangentialvektoren als *kontravariante* Vektoren, $(1, 0)$ -Tensoren, identifiziert, während Differentiale *kovariante* Vektoren, $(0, 1)$ -Tensoren, sind. Genauso wie in Kap. 3 können wir nun auch höhere Tensoren *am Punkt* P als eine multilineare Abbildung von entsprechend vielen Kopien von $T_P M$ und $T_P^* M$ nach $\mathbb{R}$ definieren. Wir sprechen also von einem (r, s) -Tensor T am Punkt P als einer Abbildung

$$T : \underbrace{T_P^* M \times \cdots \times T_P^* M}_{r \text{ Stück}} \times \underbrace{T_P M \times \cdots \times T_P M}_{s \text{ Stück}} \rightarrow \mathbb{R},$$

die linear in jedem Faktor ist. Die Koordinatendarstellung dieses Tensors bzgl. einer Basis ergibt sich in gewohnter Weise, indem man die Abbildung auf allen möglichen Kombinationen von Basis- und dualen Basisvektoren auswertet. Dies ergibt einen Satz von Zahlen

$$T^{i_1 i_2 \dots i_r}_{j_1 j_2 \dots j_s}(P),$$

deren Transformationsverhalten bei Basiswechsel durch die Zahlen (r, s) eindeutig charakterisiert ist, vgl. Kap. 3.

Nun können wir uns vorstellen, dass an jedem Punkt P der Mannigfaltigkeit M ein (r, s) -Tensor ausgewählt wird. Wir sprechen dann von einem (r, s) -**Tensorfeld**, wel- Tensorfeld

ches auf M definiert wird. Ein solches Feld von Tensoren heißt *differenzierbar*, wenn die Koordinatendarstellung des Tensors bzgl. einer natürlichen Basis, also das Schema

$$T^{i_1 i_2 \dots i_r}_{j_1 j_2 \dots j_s}(x)$$

in Abhängigkeit von den Koordinaten x^i differenzierbare Funktionen sind.

Übung 5.8: Warum ist diese Definition koordinatenunabhängig?

Ein Tensorfeld ist demnach nichts anderes als die Auswahl eines Tensors im Tangentialraum an jedem Punkt der Mannigfaltigkeit.

Einige Beispiele:

- | | |
|------------|--|
| 1-Form | (1) Mit einer Funktion f auf M ist an jedem Punkt P das totale Differential df_P definiert. Die Zuordnung $P \mapsto df_P$ definiert ein $(0, 1)$ -Tensorfeld, auch 1-Form genannt. |
| Vektorfeld | (2) Ein $(1, 0)$ -Tensorfeld, also eine Abbildung X , die jedem Punkt $P \in M$ einen Tangentialvektor $X_P \in T_P M$ zuordnet, nennt man auch Vektorfeld . |
| | (3) In jedem Tangentialraum $T_P M$ gibt es die Identitätsabbildung, den Kronecker-Tensor δ_P . Die Zuordnung $P \mapsto \delta_P$ definiert ein $(1, 1)$ -Tensorfeld. |
| Metrik | (4) Eine Abbildung g , die jedem Punkt $P \in M$ ein inneres Produkt η in $T_P M$ zuordnet, definiert ein $(0, 2)$ -Tensorfeld, auch Metrik genannt. |

Analog zum Kapitel 3 werden wir die Tensoren in $T_P M$ und die Tensorfelder auf M mit abstrakten Indizes bezeichnen, um ihren algebraischen Charakter, ihr Transformationsverhalten, sichtbar zu machen und besser rechnen zu können. So bezeichnet z. B. X^a je nach Kontext ein Vektorfeld oder einen Vektor an einem Punkt P . Ebenso sind ω_a , T^{ab} , S^a_b und δ^a_b eine 1-Form (Kovektor), ein $(2, 0)$ -Tensor(feld), ein $(1, 1)$ -Tensor(feld) bzw. der (das) Kronecker-Tensor(feld).

5.3 Der affine Zusammenhang

Wir haben nun unsere Mannigfaltigkeiten mit Koordinaten, Tangentialvektorräumen und Tensorfeldern ausgestattet. Was uns noch fehlt, ist etwas, was den Translationen in der SRT entspricht. Diese waren *das* Hilfsmittel, wenn es darum ging, zwei Ortsvektoren, die an verschiedenen Punkten angeheftet waren, miteinander in Beziehung zu setzen. Sie gestatteten uns, Vektoren von einem Punkt zu einem anderen zu transportieren.

Wir hatten gesehen, dass die Existenz der Translationen zur Folge hatte, dass die Vektorräume der Ortsvektoren alle kanonisch miteinander identifiziert werden konnten. Das Beispiel der Kugel hat uns gezeigt, dass dies für gekrümmte Mannigfaltigkeiten nicht mehr der Fall sein kann. Wir benötigen demnach etwas anderes, was aber die gleiche Funktion wie die Translationen hat, das Vergleichen von Vektoren an verschiedenen Punkten.

5.3.1 Motivation

Im Sinne des angesprochenen Vorgehens müssen wir also nun die Wirkung der Translationen *lokalisieren*. Wie haben wir uns dies vorzustellen? Zunächst sind die Translationen im Minkowski-Raum zwischen Ereignissen wirksam, die beliebig weit voneinander entfernt sein können. Außerdem wirken sie entlang einer Geraden. Wir kommen von diesen beiden Eigenschaften weg, indem wir uns zuerst die beiden Punkte, an denen wir Vektoren vergleichen wollen, durch eine Kurve verbunden denken. Dann stellen wir uns vor, dass wir die Kurve in lauter kleine Teilkurven unterteilen. Wenn wir wissen, wie sich Vektoren zwischen den Endpunkten der einzelnen Teilstücke ändern, dann können wir auch Vektoren an Anfangs- und Endpunkt der Kurve vergleichen. Im Grenzfall unendlich vieler, unendlich kurzer Teilstücke werden wir also vor die Aufgabe gestellt, Vektoren an Punkten zu vergleichen, die nur eine infinitesimal voneinander entfernt sind, also deren Änderung entlang infinitesimaler Kurvenstücke zu bestimmen. Das heißt, wir müssen einen *Ableitungsbegriff für Vektorfelder* (und, daraus abgeleitet, auch für Tensorfelder) einführen.

Wie ist das zu bewerkstelligen? Sei also $P = \gamma(\lambda)$ der Startpunkt der Kurve $\gamma(\lambda)$ und es sei $P_\lambda = \gamma(\lambda)$ ein Punkt auf der Kurve. Ein benachbarter Punkt auf der Kurve ist von der Form

$$\gamma(\lambda) \approx \gamma(0) + \dot{\gamma}(0)\lambda + \mathcal{O}(\lambda^2).$$

Das heißt, von P aus eine Distanz λ in Richtung des Tangentialvektors $\dot{\gamma}(0)$ an die Kurve verschoben.

Ein Vektorfeld X , das in der Umgebung von P definiert ist, ändert sich dann entlang des infinitesimalen Kurvenstücks gemäß

$$X(P_\lambda) \approx X(P) + „dX“ \cdot \dot{\gamma}(0) \lambda + \mathcal{O}(\lambda^2)$$

wobei hier durch die Apostrophe angedeutet sein soll, dass nicht klar ist, welche Bedeutung der Term haben sollte. Aber offensichtlich ist, dass die Änderung von X entlang der Kurve im Grenzfall gegeben ist durch einen Ausdruck der Form

$$\lim_{\lambda \rightarrow 0} \frac{X(P_\lambda) - X(P)}{\lambda} = „dX“ \cdot \dot{\gamma}(0).$$

Wir erwarten, dass dieser Ausdruck, als infinitesimale Änderung eines Vektorfeldes bei P wieder einen Tangentialvektor bei P liefert. Außerdem können wir beliebige Kurven durch P betrachten und die momentane Änderung von X entlang dieser Kurven berechnen. Das bedeutet, wir können für $\dot{\gamma}(0)$ jeden beliebigen Tangentialvektor V bei P einsetzen. Das wiederum bedeutet, dass das Objekt „dX“ jedem Tangentialvektor bei P wieder einen solchen zuordnet. Es muss sich also um eine Abbildung von $T_P M$ in sich und daher um einen $(1, 1)$ -Tensor handeln.

Wie können wir ein solches Objekt definieren? Eine erste Idee könnte darin bestehen, ein Koordinatensystem $(x^i)_{i=1:n}$ um einen Punkt P herum einzuführen, und das Vektorfeld X in der dadurch bereitgestellten natürlichen Basis auszudrücken¹

$$X = \sum_{i=1}^n X^i \frac{\partial}{\partial x^i}.$$

Dies liefert uns n Funktionen $X^i(x)$. Nun können wir die n^2 Funktionen

$$\frac{\partial X^i}{\partial x^k} \tag{5.3.1}$$

bilden. Betrachten wir ein Vektorfeld entlang einer Kurve $\gamma(\lambda)$ mit Koordinatendarstellung $x^i(\lambda)$, dann ergibt sich für die momentane Änderung der Funktionen X^i bei P

$$\frac{dX^i(\gamma(\lambda))}{d\lambda}(0) = \sum_{k=1}^n \frac{\partial X^i}{\partial x^k}(x(0)) \dot{x}^k(0),$$

also tatsächlich ein Ausdruck, der einem Tangentialvektor n Komponentenfunktionen zuordnet. Aber sind diese n Funktionen auch wirklich Komponenten eines Tensorfeldes? Oder anders ausgedrückt, sind die n^2 Funktionen (5.3.1) tatsächlich die Komponenten eines $(1, 1)$ -Tensorfeldes?

Übung 5.9: Zeigen Sie, dass sich die n^2 Funktionen (5.3.1) bei Koordinatenwechsel $y^k = \phi^k(x^i)$ wie folgt transformieren

$$\sum_{l=1}^n \frac{\partial \hat{X}^k}{\partial y^l} \phi^l_i = \sum_{l=1}^n \frac{\partial^2 \phi^k}{\partial x^i \partial x^l} X^l + \sum_{l=1}^n \phi^k_l \frac{\partial X^l}{\partial x^i}.$$

Offensichtlich können diese Funktionen nicht die Komponenten eines $(1, 1)$ -Tensorfeldes sein, denn sie transformieren sich in einer *inhomogenen* Weise.

Bemerkung. Zur Klarstellung: Wir können natürlich immer in der Karte der Koordinaten (x^i) durch die Funktionen (5.3.1) ein $(1, 1)$ -Tensorfeld T definieren, indem wir setzen

$$T = \sum_{i,k=1}^n \frac{\partial X^i}{\partial x^k} dx^k \otimes \frac{\partial}{\partial x^i}.$$

Dadurch ist offensichtlich ein Tensorfeld im Bereich der x^i -Koordinaten definiert. Beim Übergang zu einer anderen Koordinatenbasis lassen sich in üblicher Weise die Tensorkomponenten in der anderen Basis berechnen. Diese n^2 Funktionen stimmen aber nicht mit denen überein, die man bekommt, wenn man die partiellen Ableitungen der Komponenten $\hat{X}^k$ des Vektorfeldes nach den Koordinaten (y^k) berechnet.

¹Man beachte, dass hier das Argument weggelassen ist, um anzudeuten, dass es sich um ein Vektorfeld handelt, das in einer Umgebung von P definiert ist.

5.3.2 Die kovariante Ableitung

Man kann nun aufgrund dieses Transformationsverhaltens versuchen, eine befriedigende Definition für die Ableitung eines Vektorfeldes zu finden, indem man nach Wegen sucht, den inhomogenen ersten Term zum Verschwinden zu bringen. Wir wollen uns hier diesen Weg sparen (vgl. dazu Anhang A) und gleich zum Endergebnis kommen.

Wir suchen also einen **Ableitungsoperator** ∇ , der einem Vektorfeld X ein $(1, 1)$ -Tensorfeld ∇X zuordnet, die so genannte **kovariante Ableitung** von X . Wir fordern von diesem Ableitungsoperator die folgenden Eigenschaften

Ableitungsoperator
kovariante
Ableitung

- (i) ∇ ist additiv: für je zwei Vektorfelder X und Y gilt

$$\nabla(X + Y) = \nabla X + \nabla Y$$

- (ii) ∇ genügt der Produktregel: für jedes Vektorfeld X und jede Funktion f gilt

$$\nabla(fX) = df \otimes X + f \nabla X$$

Damit (ii) tatsächlich als Produktregel interpretiert werden kann, setzt man noch per Definition die Wirkung des Ableitungsoperators auf Funktionen fest als

$$\nabla f := df.$$

Bemerkung. Der Ableitungsoperator ∇ ordnet jedem Vektorfeld ein $(1, 1)$ -Tensorfeld zu, ist aber selbst kein $(1, 2)$ -Tensorfeld. Ein solches Tensorfeld, sagen wir D , würde zwar ebenfalls jedem Vektorfeld ein $(1, 1)$ -Tensorfeld zuordnen, jedoch so, dass das Bild von X am Punkt P nur vom Wert von X am Punkt P abhängig ist. Das Tensorfeld D würde also die Gleichung $D(fX) = f D(X)$ erfüllen und nicht die Produktregel (ii). Die Erklärung dafür ist, dass der Ableitungsoperator eine *Ableitung* darstellen soll, und deshalb eine momentane Änderung liefert, d.h. eine Differenz des Vektorfeldes X an *zwei* infinitesimal benachbarten Punkten.

Der Ableitungsoperator wird in Index-Notation mit ∇_a bezeichnet. Die kovariante Ableitung eines Vektorfeldes X^a wird dem entsprechend mit $\nabla_a X^b$ bezeichnet. Die Produktregel lautet dann

$$\nabla_a(fX^b) = \nabla_a f X^b + f \nabla_a X^b$$

woraus auch hervorgeht, dass wir das totale Differential einer Funktion in Index-Notation mit $\nabla_a f$ bezeichnen. Für jeden Tangentialvektor T ist die duale Paarung

$$\langle df_P, T \rangle = T^a \nabla_a f = \nabla_T f$$

die Richtungsableitung von f in Richtung des Tangentialvektors T . Das motiviert die Bezeichnung *Richtungsableitung des Vektorfeldes* X in Richtung des Tangentialvektors T für

$$T^a \nabla_a X^b \doteq \nabla_T X.$$

Es seien nun $(e_1, e_2, \dots, e_n)$ n Vektorfelder, die in einer Umgebung von P an jedem Punkt linear unabhängig sind, also in jedem Tangentialraum eine Basis bilden. Dies können z.B. die Vektoren einer natürlichen Basis sein. Wir interessieren uns für die Wirkung von ∇ auf diese Vektoren. Wir betrachten also die Richtungsableitungen

$$\nabla_{e_j} e_k$$

für je zwei Basisvektoren. Dies ist an jedem Punkt ein Vektor, kann also bezüglich der Basis dargestellt werden. So ergibt sich die Darstellung

$$\nabla_{e_j} e_k = \Gamma_{jk}^i e_i.$$

Dadurch werden n^3 Funktionen Γ_{jk}^i in der Umgebung von P definiert. Man nennt sie die **Zusammenhangskoeffizienten** von ∇ bezüglich der Basis $(e_1, e_2, \dots, e_n)$. Im Fall einer natürlichen Basis redet man auch von den **Christoffel-Symbolen** (2. Art).

Zusammenhangs-
koeffizienten
Christoffel-Symbole

Wenn wir diese Koeffizienten kennen, dann ist ein Leichtes, die kovariante Ableitung für beliebige Vektorfelder zu berechnen. Sind z.B. in einer Umgebung von P Koordinaten $(x^i)_{i=1:n}$ gegeben, dann können wir jedes Vektorfeld in dieser Umgebung bezüglich der natürlichen Basis $\partial_i := \partial/\partial x^i$ darstellen. Ist z.B. das Vektorfeld $X = \sum_{i=1}^n X^i \partial_i$ und der Tangentialvektor $T = \sum_{i=1}^n T^i \partial_i$ am Punkt P gegeben, dann berechnen wir nach den Regeln (i) und (ii)

$$\begin{aligned} \nabla_T X &= \nabla_T \left(\sum_{i=1}^n X^i \partial_i \right) = \sum_{i=1}^n \nabla_T (X^i \partial_i) \\ &= \sum_{i=1}^n (\nabla_T X^i) \partial_i + \sum_{i=1}^n X^i (\nabla_T \partial_i) \\ &= \sum_{i,k=1}^n (T^k \partial_k X^i) \partial_i + \sum_{i,k=1}^n X^i T^k (\nabla_{\partial_k} \partial_i) \\ &= \sum_{i,k=1}^n (T^k \partial_k X^i) \partial_i + \sum_{i,k=1}^n X^i T^k \left(\sum_{m=1}^n \Gamma_{ki}^m \partial_m \right) \\ &= \sum_{i,k=1}^n \left[T^k \left(\partial_k X^i + \sum_{m=1}^n \Gamma_{kl}^i X^l \right) \right] \partial_i \end{aligned}$$

Aus dieser Rechnung ergeben sich unter anderem auch die Komponenten von $\nabla_a X^b$ in einer natürlichen Basis zu

$$\partial_k X^i + \sum_{m=1}^n \Gamma_{kl}^i X^l.$$

Wir sehen, dass der erste Bestandteil dieser Komponenten die partiellen Ableitungen der Komponenten von X sind, die wir oben in Betracht gezogen hatten. Der Zusatzterm jedoch enthält die Zusammenhangskoeffizienten. Nun sind die partiellen Ableitungen

nicht die Komponenten eines Tensorfeldes, der gesamte Ausdruck aber schon. Das bedeutet also, dass auch der zweite Term für sich genommen keinen Tensor darstellen kann.

Übung 5.10: Man bestimme das Transformationsverhalten der Zusammenhangskoeffizienten unter Basiswechsel und zeige, wie sich die störenden Terme, die von den partiellen Ableitungen herrühren, gerade gegen diejenigen, die bei der Transformation der Γ 's entstehen, kompensieren.

Nun können wir zwar Vektorfelder ableiten, aber wie sollen wir andere Tensorfelder differenzieren? Ein Standardverfahren erlaubt uns, die Definition der kovarianten Ableitung auf beliebige Tensorfelder auszudehnen, wenn sie schon für Funktionen und Vektorfelder definiert ist.

Wir illustrieren das Verfahren am Beispiel von kovarianten Vektorfeldern (1-Formen). Wir betrachten ein Vektorfeld X^a und eine 1-Form ω_a auf M . Dann ist $\langle \omega, X \rangle = \omega_a X^a$ eine Funktion auf M . Wir können also ihre kovariante Ableitung, d.h., ihr totales Differential berechnen. Nun nehmen wir an, die kovariante Ableitung sei auch für 1-Formen definiert, und erfülle die Produktregel, so dass

$$\nabla_a (\omega_b X^b) = (\nabla_a \omega_b) X^b + \omega_b (\nabla_a X^b)$$

und damit auch

$$(\nabla_a \omega_b) X^b = \nabla_a (\omega_b X^b) - \omega_b (\nabla_a X^b)$$

gilt. Auf rechten Seite dieser Gleichung wirkt der Ableitungsoperator nur auf Größen, auf denen seine Wirkung bekannt ist, während er auf der linken Seite auf die 1-Form wirkt. Wir können also diese Gleichung als *Definition* der kovarianten Ableitung eines $(0, 1)$ -Tensorfelds betrachten. Diese Definition hat dann zur Folge, dass der Ableitungsoperator die oben angegebene Produktregel erfüllt.

Um zu zeigen, wie sich diese Definition auswirkt, berechnen wir die Komponenten der kovarianten Ableitung einer 1-Form ω in einer Koordinatenbasis ∂_i wie oben. Die 1-Form selbst wird dargestellt durch

$$\omega = \sum_{i=1}^n \omega_i dx^i.$$

Die Komponenten von $\nabla_a \omega_b$ erhalten wir, indem wir diesen $(0, 2)$ -Tensor auf die Basisvektoren anwenden, also berechnen wir

$$\begin{aligned} \partial_i^a \nabla_a \omega_b \partial_k^b &= \langle \nabla_{\partial_i} \omega, \partial_k \rangle = \nabla_{\partial_i} \langle \omega, \partial_k \rangle - \langle \omega, \nabla_{\partial_i} \partial_k \rangle \\ &= \partial_i \langle \omega, \partial_k \rangle - \left\langle \omega, \sum_{l=1}^n \Gamma_{ik}^l e_l \right\rangle \\ &= \partial_i \omega_k - \sum_{l=1}^n \Gamma_{ik}^l \omega_l. \end{aligned}$$

Man vergleiche diesen Ausdruck mit demjenigen für die Komponenten der kovarianten Ableitung eines Vektorfeldes.

Schließlich kann man die kovariante Ableitung auf beliebige Tensorfelder ausdehnen, indem man fordert, dass der Ableitungsoperator additiv sei und bezüglich des Tensorprodukts die Produktregel erfülle. Es gilt dann also

$$\nabla (U \otimes V) = \nabla U \otimes V + U \otimes \nabla V.$$

in Index-Notation sieht dies so aus

$$\nabla_e (U_{b,\dots}^a V_{d,\dots}^c) = (\nabla_e U_{b,\dots}^a) V_{d,\dots}^c + U_{b,\dots}^a (\nabla_e V_{d,\dots}^c).$$

Unter Benutzung dieser Rechenregeln kann man nun die Komponenten der kovarianten Ableitung für beliebige Tensorfelder berechnen. Als Beispiel sei hier nur die Formel für ein $(1, 2)$ -Tensorfeld T_{bc}^a angegeben. Wir bezeichnen die Komponenten von $\nabla_e T_{bc}^a$ mit $T_{jk;l}^i$ und erhalten

$$T_{jk;l}^i = \partial_l T_{jk}^i + \Gamma_{lm}^i T_{jk}^m - \Gamma_{lj}^m T_{mk}^i - \Gamma_{lk}^m T_{jm}^i.$$

Übung 5.11: Verifizieren Sie diese Formel.

Zum Schluss dieses Abschnitts müssen wir uns fragen, ob es auf einer beliebigen n -dimensionalen Mannigfaltigkeit überhaupt einen Ableitungsoperator gibt. Diese Frage lässt sich positiv beantworten: Es gibt immer einen Ableitungsoperator. Aber damit nicht genug, denn es gibt sogar beliebig viele.

Ist mit ∇ auch $\hat{\nabla}$ ein Ableitungsoperator, dann gilt

$$\hat{\nabla} f - \nabla f = 0$$

für jede Funktion f . Außerdem ist durch

$$(\omega, X, Y) \mapsto \langle \omega, \hat{\nabla}_X Y - \nabla_X Y \rangle = C(\omega, X, Y)$$

eine Abbildung definiert, die an jedem Punkt P dem Paar (ω_P, X_P) eine Zahl zuordnet. Sie ist offensichtlich linear in ω_P und X_P . Ist sie auch linear in Y_P ? Dies ist nicht ganz offensichtlich, weil hier *Ableitungen* von Y vorkommen. Wir betrachten also den Ausdruck $C(\omega, X, fY)$ und verwenden die Rechenregeln für Ableitungsoperatoren. Dies ergibt

$$\begin{aligned} C(\omega, X, fY) &= \langle \omega, \hat{\nabla}_X (fY) - \nabla_X (fY) \rangle \\ &= \langle \omega, X(f)Y + f \hat{\nabla}_X Y - X(f)Y - f \nabla_X Y \rangle \\ &= fC(\omega, X, Y). \end{aligned}$$

An jedem Punkt P gilt also $C(\omega, X, fY)_P = f(P)C(\omega, X, Y)_P$. Damit ist $C(\omega, X, Y)_P$ auch linear in Y , d.h. es definiert ein $(1, 2)$ -Tensorfeld $C_a{}^b{}_c$. Wir erhalten also die Gleichung

$$\hat{\nabla}_a X^b = \nabla_a X^b + C_a{}^b{}_c X^c.$$

Hier ist $C_a{}^b{}_c$ durch die Differenz der Ableitungsoperatoren bestimmt. Aber wir können umgekehrt feststellen, dass für einen beliebigen Tensor $C_a{}^b{}_c$ durch obige Gleichung ein weiterer Ableitungsoperator definiert wird.

Übung 5.12: Zeigen Sie, dass der Ableitungsoperator $\hat{\nabla}$, der dadurch definiert wird, die Rechenregeln für Ableitungsoperatoren erfüllt.

Man hat also im Allgemeinen die Qual der Wahl eines Ableitungsoperators. Wir werden später sehen, dass sich durch Einführung einer Metrik auf M diese Wahl drastisch reduziert: es gibt dann nämlich nur noch einen geeigneten Ableitungsoperator.

5.3.3 Parallel-Transport

Wie löst nun der Ableitungsoperator das Problem des Vergleichs von Vektoren? Zunächst betrachten wir die Situation der infinitesimalen Änderung eines Vektorfeldes X^a entlang einer Kurve $\gamma(\lambda)$ mit Tangentialvektor T^a durch $P \in M$. Dann ist die kovariante Ableitung von X^a entlang der Kurve bei P gegeben durch den Vektor $\nabla_T X$, der in Index-Notation als

$$T^a \nabla_a X^b,$$

geschrieben wird. D.h. man berechnet ihn, indem man $\nabla_a X^b$ mit T^a überschreibt. Wenn wir also den Ableitungsoperator ∇_a , also seine Zusammenhangskoeffizienten Γ_{jk}^i , kennen, dann können wir diesen Vektor berechnen. Der Ableitungsoperator definiert also, welche infinitesimale Änderung das Vektorfeld X entlang der Kurve bei P erfährt.

Betrachten wir nun eine Kurve γ zwischen zwei verschiedenen Punkten P und Q auf M . Auf dieser Kurve sei ein Vektorfeld X gegeben. An jedem Punkt der Kurve können wir demnach die infinitesimale Änderung von X entlang des dortigen Tangentialvektors T berechnen. Dies liefert uns ein weiteres Vektorfeld $\nabla_T X$ auf γ . Man trifft nun die folgende

Definition 5.2. Ein Vektorfeld X heisst parallel entlang einer Kurve γ zwischen zwei Punkten P und Q , wenn an jedem Punkt der Kurve

$$T^a \nabla_a X^b = 0$$

gilt. Man bezeichnet das Vektorfeld auch als parallel von P nach Q verschoben.

Diese Sprechweise kommt daher, dass man entlang der Kurve von P nach Q bei gegebenem Vektor X_P am Punkt P ein eindeutiges Vektorfeld X auf der Kurve finden kann, welches parallel ist und bei P mit X_P übereinstimmt. Man kann also dieses Vektorfeld als die Zwischenstadien betrachten, die der Vektor X_P durchläuft, wenn man ihn so

von P nach Q bringt, dass er keine infinitesimale Änderung erfährt, also immer konstant bleibt. Dieser sogenannte **Paralleltransport** ist im Allgemeinen von der Kurve abhängig, vgl. Abb. 5.6

Paralleltran

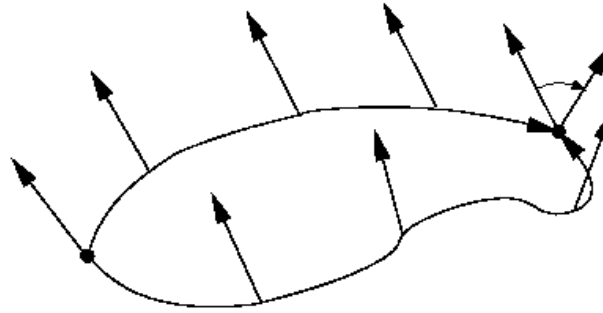


Abbildung 5.6: Der Paralleltransport eines Vektors ist wegabhängig

Bemerkung. Ebenso wie man von parallelen Vektorfeldern spricht, kann man auch parallele Tensorfelder betrachten. Das sind Tensorfelder $U^a_{b,\dots}$, definiert entlang einer Kurve, deren infinitesimale Änderung an jedem Kurvenpunkt verschwindet. Für sie gilt also die Gleichung

$$\nabla_T U = 0 \doteq T^c \nabla_c U^a_{b,\dots} = 0.$$

Übung 5.13: Man zeige, dass die Gleichung $T^a \nabla_a Z^b = 0$ in Koordinaten die Darstellung

$$\dot{Z}^i(\lambda) + \sum_{j,k=1}^n \Gamma^i(x(\lambda))_{jk} \dot{x}^j(\lambda) Z^k(\lambda) = 0$$

besitzt. Dabei sind die $Z^i(\lambda) = Z^i(x(\lambda))$ die Komponenten des Vektorfelds Z auf der Kurve mit Darstellung $x^i(\lambda)$.

Ein besonderes Vektorfeld, welches entlang jeder Kurve definiert ist, ist das Tangentialvektorfeld, das an jedem Punkt der Kurve den Tangentialvektor der Kurve spezifiziert. Wenn man nun verlangt, dass dieses Vektorfeld parallel sein soll, so ist dies nicht nur eine Bedingung an das Vektorfeld, sondern viel mehr noch an die Kurve selbst. Man nennt eine Kurve, deren Tangentialvektor parallel entlang der Kurve selbst verschoben ist, eine **Autoparallele**. Die Kurve ist daher so beschaffen, dass der Tangentialvektor immer parallel zu sich selbst verschoben wird². Eine Autoparallele γ mit Tangentialvektorfeld T genügt folglich der Gleichung

Autoparallele

$$T^a \nabla_a T^b = 0 \doteq \nabla_T T = 0.$$

²Der Begriff Autoparallele wird manchmal etwas weiter gefasst. Ist die infinitesimale Änderung des Tangentialvektors der Kurve proportional zum Tangentialvektor, dann ändert sich seine Richtung nicht. Deshalb werden auch Kurven deren Tangentialvektorfeld die Gleichung $\nabla_T T \propto T$ erfüllt, auch als Autoparallele bezeichnet.

Übung 5.14: Zeigen Sie, dass die Koordinatendarstellung $x^i(\lambda)$ einer Autoparallelen der Gleichung

$$\ddot{x}^i(\lambda) + \Gamma_{jk}^i(x(\lambda))\dot{x}^j(\lambda)\dot{x}^k(\lambda) = 0$$

genügt.

Eine autoparallele Kurve ist die Verallgemeinerung einer Geraden auf einer gekrümmten Mannigfaltigkeit, eine Kurve deren Tangentialvektor nie die Richtung ändert.

5.4 Krümmung

Der Paralleltransport eines Vektors vom Punkt P zum Punkt Q ist im Allgemeinen wegababhängig. Das heisst (vgl. Abb. 5.7), der Paralleltransport eines Vektors X von P über P' zu

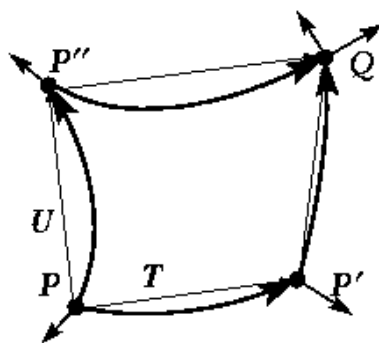


Abbildung 5.7: Zur Nichtkommutativität der kovarianten Ableitung

zu Q wird im Allgemeinen ein anderes Ergebnis liefern, als der Paralleltransport über P''. Wie kann man den Unterschied dieser Vektoren fassen?

Im Sinne des üblichen Vorgehens werden wir versuchen, die Situation zu lokalisieren, das heisst, sie durch infinitesimale Änderungen zu beschreiben. Stellen wir uns vor, der Punkt Q (und damit die Punkte P' und P'') rückt immer näher auf den Punkt P zu. Dann können wir uns anschaulich leicht davon überzeugen, dass der Unterschied des Paralleltransports von X herrührt von der Vertauschung der kovarianten Ableitungen in Richtung T bzw. U. Dies motiviert die Untersuchung von

$$\nabla_T \nabla_U X - \nabla_U \nabla_T X$$

bzw., wenn wir die Betrachtung für alle möglichen Kurvenpaare durchführen, des Kommutators der kovarianten Ableitungen

$$\Delta_{ab} = \nabla_a \nabla_b - \nabla_b \nabla_a.$$

Wir untersuchen zuerst die Wirkung dieses Operators auf eine Funktion f. Offensichtlich erhalten wir das (0, 2)-Tensorfeld $\Delta_{ab}f$. Wir betrachten die Wirkung des Kommutators auf einem Produkt von Funktionen. Es ist leicht zu verifizieren, dass

$$\Delta_{ab}(fg) = (\Delta_{ab}f)g + f(\Delta_{ab}g)$$

gilt. Das bedeutet aber, dass für beliebige Vektoren X^a und Y^b gilt, dass

$$T(X, Y) = X^a Y^b \Delta_{ab}$$

eine Derivation auf Funktionen ist, also einen Tangentialvektor definiert. Das heisst, es gibt ein $(1, 2)$ -Tensorfeld $T_{ab}{}^c$, für das

$$T(X, Y) = T_{ab}{}^c X^a Y^b$$

Torsion gilt. Dieses Tensorfeld nennt man die **Torsion** des Zusammenhangs bzw. des Ableitungsoperators. Es gilt also die Gleichung

$$\nabla_a \nabla_b f - \nabla_b \nabla_a f = T_{ab}{}^c \nabla_c f$$

für beliebige Funktionen f . Der Torsionstensor misst also die Nichtkommutativität der kovarianten Ableitung *in ihrer Wirkung auf Funktionen*. Nun hatten wir ja gesehen, dass wir viele Möglichkeiten bei der Auswahl eines Ableitungsoperators ∇_a haben. Diese Freiheit in der Wahl des Zusammenhangs kann man nun einschränken, indem man verlangt, dass der Zusammenhang *torsionsfrei* sei.

Übung 5.15: Man beweise, dass es immer möglich ist, ausgehend von einem beliebigen Ableitungsoperator $\hat{\nabla}_a$ einen torsionsfreien Ableitungsoperator ∇_a einzuführen.

Übung 5.16: Bestimmen Sie die Tensorkomponenten des Torsionstensors $T_{ab}{}^c$ und zeigen Sie damit, dass die Koeffizienten Γ_{jk}^i eines torsionsfreien Zusammenhangs symmetrisch sind

$$\Gamma_{jk}^i = \Gamma_{kj}^i.$$

Wir können und werden uns also in Zukunft auf torsionsfreie Zusammenhänge beschränken, also immer annehmen, dass

$$\nabla_a \nabla_b f = \nabla_b \nabla_a f$$

gilt.

Betrachten wir nun die Wirkung von Δ_{ab} auf Vektorfelder X^d . Wir untersuchen den Ausdruck

$$\Delta_{ab} X^d.$$

Es sieht zunächst so aus, als ob dieser Ausdruck von der zweiten Ableitung des Vektorfeldes abhängen würde. Wir zeigen nun aber, dass dies nicht so ist, sondern dass er tatsächlich nur vom Wert von X selber abhängig ist. Dazu betrachten wir

$$\Delta_{ab} (f X^d) = (\Delta_{ab} f) X^d + f (\Delta_{ab} X^d) = f (\Delta_{ab} X^d).$$

Das heisst, die Abbildung

$$(X, Y, Z) \mapsto R(Y, Z)X \doteq (X^d, Y^a, Z^b) \mapsto Y^a Z^b \Delta_{ab} X^d$$

die den drei Vektorfeldern wieder ein Vektorfeld zuordnet, liefert an jedem Punkt P der Mannigfaltigkeit eine multilineare lineare Abbildung $R_P : T_P M \times T_P M \times T_P M \rightarrow T_P M$, also einen $(1, 3)$ -Tensor in $T_P M$. Insgesamt wird daher ein $(1, 3)$ -Tensorfeld

$$R_{abc}{}^d$$

Gleichung definiert. Für dieses Tensorfeld gilt die **Ricci-Gleichung**

$$(\nabla_a \nabla_b - \nabla_b \nabla_a) X^d = R_{abc}{}^d X^c$$

Das Tensorfeld $R_{abc}{}^d$ ist der sogenannte **Riemann-Tensor**. Er misst die Nichtkommutativität der kovarianten Ableitung in ihrer Wirkung auf Vektoren und ist damit auch ein Maß für die Nichtkommutativität des Paralleltransports entlang verschiedener Wege.

Riemann-Tensor

Die Eigenschaften des Riemann-Tensors folgen sofort aus seiner Definition.

- (i) Zunächst gilt offensichtlich die Symmetrie-Eigenschaft

$$R_{abc}{}^d = R_{bac}{}^d \iff R_{(ab)c}{}^d = 0.$$

- (ii) Der Riemann-Tensor erfüllt die **zyklische Identität**, auch erste Bianchi-Identität genannt:

zyklische Identität

$$R_{[abc]}{}^d = 0 \iff R_{abc}{}^d + R_{bca}{}^d + R_{cab}{}^d = 0.$$

- (iii) Der Riemann-Tensor erfüllt die (zweite) **Bianchi-Identität**

Bianchi-Identität

$$\nabla_{[e} R_{ab]c}{}^d = 0$$

Übung 5.17: Zeigen Sie, dass der Kommutator Δ_{ab} der kovarianten Ableitungen angewandt auf ein Tensorprodukt die Produktregel erfüllt, so dass also

$$\Delta_{pq} (U_{b\dots}^a \dots V_{d\dots}^c \dots) = (\Delta_{pq} U_{b\dots}^a \dots) V_{d\dots}^c \dots + U_{b\dots}^a \dots (\Delta_{pq} V_{d\dots}^c \dots).$$

gilt.

Übung 5.18: Zeigen Sie die zyklische Identität für den Riemann-Tensor, indem Sie die total antisymmetrisierte dreifache kovariante Ableitung einer Funktion f

$$\nabla_{[a} \nabla_b \nabla_{c]} f$$

auf zwei verschiedene Arten auswerten.

Übung 5.19: Zeigen Sie auf ähnliche Weise die zweite Bianchi-Identität, indem Sie die gleiche total antisymmetrisierte dreifache kovariante Ableitung eines Vektorfeldes X

$$\nabla_{[a} \nabla_b \nabla_{c]} X^d$$

auf zwei verschiedene Arten auswerten.

Der Riemann-Tensor lässt sich einmal kontrahieren zum **Ricci-Tensor**

Ricci-Tensor

$$R_{ab} = R_{aeb}{}^e.$$

5.5 Riemannsche Geometrie

Die Beschränkung auf torsionsfreie Zusammenhänge läßt immer noch eine riesige Auswahl an Ableitungsoperatoren offen. Um sie weiter einschränken zu können, erinnern wir uns an die Parallelverschiebung im euklidischen Raum. Neben vielen anderen, hat sie die Eigenschaft, dass sie Strecken und Winkel invariant läßt. So gehen zum Beispiel Dreiecke durch Parallelverschiebung in kongruente Dreiecke über. Strecken und Winkel werden durch eine Metrik definiert, im euklidischen Raum also eine positiv definite Bilinearform. Im Rahmen der SRT hatten wir die Raumzeit-Metrik ebenfalls als symmetrische Bilinearform kennen gelernt, die zwar nicht positiv definit, aber doch nicht entartet war. Die entsprechende Verallgemeinerung auf Mannigfaltigkeiten ist die

Metrik **Definition 5.3.** Eine **Metrik** g auf einer Mannigfaltigkeit M ist ein symmetrisches, reguläres $(0, 2)$ -Tensorfeld. Ist insbesondere M vier-dimensional und hat g die Signatur $(+, -, -, -)$, so sprechen wir von einer **Lorentz-Metrik**.

Eine solche Metrik erzeugt an jedem Punkt P der Mannigfaltigkeit eine nicht entartete, symmetrische Bilinearform im Tangentialraum der Mannigfaltigkeit in P . Dies ist in genauer Analogie zum euklidischen Raum bzw. Minkowski-Raum zu sehen, wo die Metrik ebenfalls eine Bilinearform für die Vektoren *an jedem Punkt* war. Die Tatsache, dass sich die Metrik nicht von Punkt zu Punkt ändert, liegt daran, dass die Parallelverschiebung, die uns von Punkt zu Punkt bringt, eben diese Metrik invariant läßt.

Dies motiviert uns, die kovariante Ableitung der Metrik für einen gegebenen torsionsfreien Ableitungsoperator zu betrachten. Dies ist das Tensorfeld

$$\Delta_{abc} = \nabla_a g_{bc}.$$

Wir fragen nun: Ist es möglich, einen *anderen* torsionsfreien Ableitungsoperator $\tilde{\nabla}$ zu finden, so dass die Metrik für diesen Ableitungsoperator konstant ist, so dass also

$$\tilde{\nabla}_a g_{bc} = 0$$

gilt. Offensichtlich muss dann für die Differenz der Operatoren die Gleichung

$$0 = \tilde{\nabla}_a g_{bc} = \nabla_a g_{bc} - C_a^e{}_b g_{ec} - C_a^e{}_c g_{be} = \Delta_{abc} - C_{acb} - C_{abc}$$

gelten. Wenn sich diese Gleichung für beliebige Δ_{abc} nach C_{abc} auflösen lässt, dann haben wir gezeigt, dass man Ableitungsoperatoren finden kann, für die die Metrik kovariant konstant ist. Zur Lösung der Gleichung schreiben wir sie dreimal mit zyklisch vertauschten Indizes auf

$$\begin{aligned} \Delta_{abc} &= C_{abc} + C_{acb}, \\ - | \Delta_{bca} &= C_{bca} + C_{bac}, \\ + | \Delta_{cab} &= C_{cab} + C_{cab} \end{aligned}$$

und erhalten nach Subtraktion der zweiten und Addition der dritten Zeile die Gleichung

$$\Delta_{abc} + \Delta_{cab} - \Delta_{bca} = (C_{abc} + C_{cba}) + (C_{acb} - C_{bca}) + (C_{cab} - C_{bac})$$

Beachten wir nun, dass $C_{abc} = C_{cba}$ gilt, da wir torsionsfreie Ableitungsoperatoren betrachten, so folgt

$$C_{abc} = \frac{1}{2} (\Delta_{abc} + \Delta_{cab} - \Delta_{bca}).$$

Diese Gleichung bedeutet, dass sich der Differenztensor durch die Forderung, dass die Metrik bzgl. des neuen Zusammenhangs konstant sei, *eindeutig* aus Δ_{abc} berechnen lässt. Daraus folgt schließlich der

Satz 5.2. *Ist auf einer Mannigfaltigkeit eine Metrik definiert, so gibt es genau einen torsionsfreien Zusammenhang, so dass die Metrik kovariant konstant ist. Dieser heißt **Levi-Civita Zusammenhang** der Metrik oder auch metrischer Zusammenhang.*

Levi-Civita
Zusammenhang

Man sagt auch, Metrik und Ableitungsoperator seien kompatibel. Dieser Satz erleichtert uns die Wahl eines Ableitungsoperators ungemein. Haben wir eine Metrik gegeben so wie im Falle der ART, dann ist uns damit auch ein eindeutiger Zusammenhang gegeben.

Da die Metrik den Zusammenhang eindeutig bestimmt, muss es möglich sein, die Zusammenhangskoeffizienten aus den Koeffizienten der Metrik zu berechnen. Sei uns also eine Metrik g_{ab} gegeben. In einer Karte mit den Koordinaten x^i sind die metrischen Koeffizienten gegeben durch die Funktionen

$$g_{ik} = g_{ab} \partial_i^a \partial_k^b = g(\partial_i, \partial_k),$$

so dass die Metrik sich in der Form

$$g = \sum_{ik} g_{ik}(x) dx^i \otimes dx^k$$

darstellen lässt. Mit Hilfe von (A.3.4) bestimmen wir die kovariante Ableitung von g und erhalten, nachdem wir die Komponenten zu Null gesetzt haben, die Gleichungen

$$\frac{\partial g_{ik}}{\partial x^l} = \Gamma_{kli} + \Gamma_{ilk}$$

wobei wir zur Abkürzung $\Gamma_{ilk} = \sum_m g_{im} \Gamma_{lk}^m$ setzen. Zur Bestimmung der Γ_{ilk} schreiben wir diese Gleichung dreimal mit zyklischer Indexpermutation auf

$$\begin{aligned} \frac{\partial g_{ik}}{\partial x^l} &= \Gamma_{kli} + \Gamma_{ilk} \\ \frac{\partial g_{kl}}{\partial x^i} &= \Gamma_{lik} + \Gamma_{kil} \\ \frac{\partial g_{li}}{\partial x^k} &= \Gamma_{ikl} + \Gamma_{lki} \end{aligned}$$

und subtrahieren die dritte von der Summe der ersten beiden. Dies ergibt mit Beachtung von $\Gamma_{ik}^m = \Gamma_{ki}^m$ wegen der Torsionsfreiheit

$$2\Gamma_{kli} = \frac{\partial g_{ik}}{\partial x^l} + \frac{\partial g_{kl}}{\partial x^i} - \frac{\partial g_{li}}{\partial x^k}.$$

Daraus ergeben sich schließlich nach Multiplikation mit der zu g_{ik} inversen Matrix g^{il} die Zusammenhangskoeffizienten des Levi-Civita Zusammenhangs

$$\Gamma_{jk}^i = \frac{1}{2} \sum_l g^{il} \left(\frac{\partial g_{jl}}{\partial x^k} + \frac{\partial g_{kl}}{\partial x^j} - \frac{\partial g_{jk}}{\partial x^l} \right).$$

Christoffel-Symbole

Die so aus der Metrik bestimmten Zusammenhangskoeffizienten nennt man auch die **Christoffel-Symbole** (2. Art)³.

Der Riemann-Tensor des metrischen Zusammenhangs besitzt zusätzlich zu den schon bekannten Eigenschaften noch eine weitere Symmetrie die sich aus der Konstanz der Metrik ergibt. Es gilt nämlich

$$0 = (\nabla_a \nabla_b - \nabla_b \nabla_a) g_{cd} = -R_{abc}{}^e g_{ed} - R_{abd}{}^e g_{ce} = R_{abcd} + R_{abdc}$$

d.h., das Tensorfeld R_{abcd} , welches aus dem Riemann-Tensor durch Senken des letzten Index gewonnen wird ist antisymmetrisch in den hinteren Indizes. Zusätzlich besitzt dieses Tensorfeld aber auch die Symmetrien des Riemann-Tensors $R_{[abc]d} = 0$ und $R_{abcd} = -R_{bacd}$.

Übung 5.20: Leiten Sie aus diesen Symmetrien die sog. Paarsymmetrie

$$R_{abcd} = R_{cdab}$$

her.

³Die Christoffel-Symbole 1. Art sind die oben eingeführten Γ_{ijk} , die man aus den Christoffel-Symbolen durch Kontraktion mit der Metrik bekommt.

6 Die Feldgleichungen der Gravitation

6.1 Motivation

Wir wiederholen noch einmal die Folgerungen aus unserer Diskussion der SRT und des Äquivalenzprinzips aus Kap. 4 und 5

1. Die Raumzeit wird als eine *4-dimensionale Mannigfaltigkeit* $\mathcal{M}$ beschrieben.
2. Sie trägt eine *Lorentz-Metrik*, also ein $(0, 2)$ -Tensorfeld g_{ab} mit Signatur $(+, -, -, -)$.
3. Kräftefreie Körper bewegen sich auf *Geodäten*.

Ein paar Worte zur Erläuterung: Punkt 1. ist leicht einzusehen. Wir können Ereignisse eindeutig charakterisieren durch die Angabe von vier reellen Zahlen. Also gibt es eine Abbildung, die jedem Ereignis vier Zahlen zuordnet. Wir sind weiterhin in der Wahl dieser vier Zahlen nicht eingeschränkt. Wir können also einer Menge von Ereignissen mehrere 'Zahlensätze' – Koordinaten – zuordnen. Offensichtlich müssen sich diese Koordinaten eineindeutig ineinander transformieren lassen damit die ganze Beschreibung sinnvoll bleibt. Nehmen wir nun noch ein paar technische Forderungen mit (Stetigkeit, Differenzierbarkeit), dann sind wir sofort bei einer Mannigfaltigkeit.

In der SRT hatten wir gesehen, dass sich an jedem Ereignis P ein *Nullkegel* definieren lässt und dass auch ein Begriff von Einheitslänge oder -dauer existiert. Dies war immer in Bezug auf Vektoren an dem betrachteten Ereignis P geschehen. Wir hatten gesehen, dass diese Eigenschaften dazu führten, eine Raumzeit-Metrik η_{ab} einzuführen, die es gestattet ein invariantes Skalarprodukt zwischen je zwei Vektoren zu berechnen.

Um diese Eigenschaften einer Raumzeit in den allgemeineren Kontext zu übernehmen, beachten wir die Tatsache, dass an jedem Punkt unserer Raumzeit-Mannigfaltigkeit ein Tangentialraum existiert, also ein 4-dimensionaler Vektorraum. Fordert man nun, dass all diese Eigenschaften aus der SRT¹ im *Tangentialraum* eines jeden Punktes gelten, dann hat man ein $(0, 2)$ -Tensorfeld wie unter Punkt 2. eingeführt.

Nun sitzt zwar an jedem Ereignis zunächst eine 'andere' Metrik, weil ja die Tangentialräume an verschiedenen Punkten nichts miteinander zu tun haben. Hier kommt aber

¹Dass man hier die Eigenschaften der SRT lokalisiert, liegt daran, dass diese Theorie experimentell sehr gut gesichert ist. Man kann ebenso die Newtonsche Raumzeit lokalisieren. Dann treten an die Stelle der Nullkegel die 'Flächen der Gleichzeitigkeit'. Ein analoges Vorgehen wie im weiteren dargestellt führt auf eine Verallgemeinerung der Newtonschen Gravitationstheorie.

jetzt ein Zusammenhang zu Hilfe. Wir hatten gesehen, dass es zu einer gegebenen Metrik genau einen Zusammenhang, den *metrischen Zusammenhang* gibt. Dieser ist uns also mit der Existenz der Metrik auch sofort in die Hände gelegt. Dieser Zusammenhang hat die erfreuliche Eigenschaft, dass die kovariante Ableitung der Metrik verschwindet. Das bedeutet, dass die Metrik bei Paralleltransport konstant bleibt. Daher sitzt in einem gewissen Sinn in jedem Ereignis ‘die gleiche Metrik’.

An jedem Ereignis P lässt sich eine Basis im Tangentialraum $T_P\mathcal{M}$ finden, bzgl. der die Metrik die Diagonalform

$$\begin{pmatrix} +1 & & & \\ & -1 & & \\ & & -1 & \\ & & & -1 \end{pmatrix}$$

besitzt oder, anders ausgedrückt, es gibt an jedem Ereignis eine ONB. Diese besteht aus einem zeitartigen Einheitsvektor und drei auf einander senkrechten raumartigen Einheitsvektoren. Dies ist gerade ein lokales Bezugssystem, ein lokales Inertialsystem.

Kräftefreie Körper laufen auf Geodäten, also den geraden Linien in der Raum-Zeit. Insbesondere sind Körper, die im Gravitationsfeld der Erde frei fallen, kräftefrei. Wie ist es zu erklären, dass wir trotzdem den Eindruck haben, diese Körper seien beschleunigt? Offensichtlich liegt dies daran, dass wir selbst beschleunigt sind. Die Erdoberfläche übt eine Kraft auf uns aus, die uns auf der Oberfläche festhält und damit vom ‘Pfad der Kräftefreiheit’, der geradlinigen Bewegung in Richtung Erdmittelpunkt, abhält. Die Beschleunigung, die wir einem fallenden Stein zumessen, ist tatsächlich unsere eigene Beschleunigung, die uns auf der Erdoberfläche ‘in Ruhe’ hält. Wir sehen also, was tatsächlich wesentlich ist, ist *die Relation zwischen* der Weltlinie des Steins und unserer eigenen.

Genauso verhält es sich in dem Beispiel eines auf die Erde zufallenden Schwarms von Körpern (vgl. Fig. 2.1). Jeder dieser Körper fällt auf einer Geodäten; jeder dieser Körper kann durch Transformation ins lokal mitbewegte Inertialsystem zur Ruhe gebracht werden. Trotzdem führen die restlichen Körper eine beschleunigte Bewegung in Bezug auf den ausgewählten Körper aus. Der Eindruck der Beschleunigung kommt also nicht durch Betrachtung eines *einzelnen* Körpers, einer einzelnen Weltlinie, zustande, sondern nur durch Betrachtung *mehrerer* Weltlinien. Deren Bewegung *relativ zueinander* erzeugt den Eindruck der Beschleunigung.

Die geraden Linien, die Geodäten, sind in Bezug auf den Zusammenhang definiert. Dieser wiederum liegt fest, wenn die Metrik gegeben ist. Wenn also die Raum-Zeit-Metrik fest liegt, dann auch die Geodäten. Das heißt, *die Metrik bestimmt das Verhalten der freien Körper*. Um dies noch etwas genauer zu fassen, betrachten wir zunächst noch einmal die Bewegung eines Körpers im Gravitationsfeld der Newtonschen Theorie.

In kartesischen Koordinaten ist die Bewegung eines Körpers im Gravitationspotential ϕ durch die Bewegungsgleichung

$$\frac{d^2 x^i}{dt^2} = -\frac{\partial \phi(x(t))}{\partial x^i}$$

gegeben. In dieser Gleichung sind die Orts- und die Zeitkoordinate nicht gleich behandelt, die ersteren sind abhängige Variablen, die Zeitkoordinate ist der Kurvenparameter. Führen wir der Gerechtigkeit halber einen anderen Parameter λ entlang der Kurve ein, so ergeben sich nacheinander

$$\begin{aligned} \frac{dx^i}{d\lambda} &= \frac{dx^i}{dt} \frac{dt}{d\lambda}, \\ \frac{d^2 x^i}{d\lambda^2} &= \frac{d^2 x^i}{dt^2} \left(\frac{dt}{d\lambda} \right)^2 + \frac{dx^i}{dt} \frac{d^2 t}{d\lambda^2} = -\frac{\partial \phi(x(t))}{\partial x^i} \left(\frac{dt}{d\lambda} \right)^2 + \frac{dx^i}{dt} \frac{d^2 t}{d\lambda^2}. \end{aligned}$$

Wir wählen nun den Parameter so, dass $\frac{d^2 t}{d\lambda^2} = 0$ ist. Dann haben wir die Bewegungsgleichung in der Form (der Punkt bedeutet Ableitung nach λ)

$$\begin{aligned} \ddot{x}^0 &= \ddot{t} = 0, \\ \ddot{x}^i + \frac{\partial \phi(x(t))}{\partial x^i} \dot{t}^2 &= 0. \end{aligned}$$

Man vergleiche dies mit der Geodätengleichung in einem Koordinatensystem

$$\ddot{x}^i + \sum_{kl} \Gamma_{kl}^i \dot{x}^k \dot{x}^l = 0, \quad i = 0 : 3.$$

Die Ähnlichkeit ist offensichtlich. Setzen wir jetzt noch die Christoffel-Symbole für den metrischen Zusammenhang ein, so ergibt sich

$$\ddot{x}^i + \frac{1}{2} \sum_{klm} g^{im} \left(\frac{\partial g_{lm}}{\partial x^k} + \frac{\partial g_{km}}{\partial x^l} - \frac{\partial g_{kl}}{\partial x^m} \right) \dot{x}^k \dot{x}^l = 0, \quad i = 0 : 3.$$

Der Vergleich der beiden Gleichungssysteme zeigt also, dass das Gravitationsfeld mit den Zusammenhangskoeffizienten in Beziehung gesetzt werden muss, während die Koeffizienten g_{ik} der Metrik mit dem Gravitationspotential eingestuft werden müssen. Während die Newtonsche Theorie mit einer einzigen Funktion ϕ auskommt, sind es in der ART zehn Potentiale. Die Newtonsche Theorie ist eine *skalare* Theorie, die ART eine *tensorielle* Theorie.

Mit diesem Ergebnis ist offensichtlich, dass die Metrik bestimmt, wie Bewegungen in der Raum-Zeit ablaufen. Wie aber wird die Metrik bestimmt? Es sollte nicht so sein, dass uns die Raum-Zeit *a priori* gegeben ist, dass sie eine Kulisse darstellt, in der sich die Physik abspielt, ohne dass die Raum-Zeit selbst daran beteiligt wäre. Dies wäre analog zur Existenz eines absoluten Raumes und einer absoluten Zeit. Wir müssen demnach Gesetzmäßigkeiten auffinden, denen sich die Geometrie der Raum-Zeit unterwirft.

Die Raum-Zeit sagt der Materie, wie sie sich bewegt. Wer aber sagt der Raum-Zeit, wie sie sich krümmen muss? Dies kann nur die Materie selbst sein. Um zu sehen, welcher Art die Gesetze sein müssen, die wir aufzustellen haben, betrachten wir wieder den Schwarm von freien Körpern im Gravitationsfeld und zwar zunächst wieder in der Newtonschen Theorie.

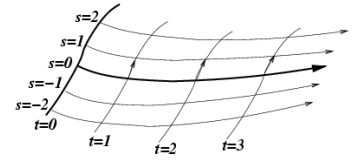


Abbildung 6.1:

Sei $x^i(t, s)$ eine 1-parametrische Schar von Teilchen, die alle zur Zeit $t = 0$ auf einer gegebenen Kurve loslaufen (vgl. Abb. 6.1). Die Startpunkte werden durch den Parameter s gegeben, welcher demnach die ganze Kurve festlegt. Die Kurven $t = \text{const.}$ sind die Punkte, die auf den Bahnen nach der gegebenen Zeit t erreicht werden. Die Menge aller dieser Bahnen bildet eine 2-dimensionale Fläche S , auf der jeder Punkt eindeutig durch die Werte der beiden Parameter s (auf welcher Bahn liegt der Punkt?) und t (zu welcher Zeit wird er erreicht?) charakterisiert werden kann.

Sei nun $z^i(t, s) = x^i(t, s) - x^i(t, 0)$ der Abstand zwischen der 'zentralen' Bahn mit $s = 0$ und einer benachbarten Bahn. Dann erfüllt dieser Abstand die Gleichung

$$\ddot{z}^i(t, s) = - \left(\frac{\partial \phi}{\partial x^i}(x(t, 0) + z(t, s)) - \frac{\partial \phi}{\partial x^i}(x(t, 0)) \right). \quad (6.1.1)$$

Für kleine Werte von s ist dies

$$\ddot{z}^i(t, s) \approx \ddot{z}^i(t, 0) + s \frac{\partial \ddot{z}^i}{\partial s}(t, 0) + \mathcal{O}(s^2) = -s \sum_k \frac{\partial^2 \phi}{\partial x^i \partial x^k}(x(t, 0)) \frac{\partial z^k}{\partial s}(t, 0) + \mathcal{O}(s^2).$$

Nach Vertauschen der Ableitungen $\partial/\partial s$ und $\partial/\partial t$ und der Definition

$$Z^i(t) = \frac{\partial z^i}{\partial s}(t, 0)$$

erhalten wir schließlich in 1. Ordnung die Gleichung

$$\ddot{Z}^i(t) = - \sum_k \frac{\partial^2 \phi}{\partial x^i \partial x^k}(x(t, 0)) Z^k(t). \quad (6.1.2)$$

Wie haben wir dieses Ergebnis zu interpretieren? Beginnen wir mit (6.1.1), wo die Interpretation noch klar ist. Diese Gleichung regelt den Abstand der Teilchen auf der Bahn $s = 0$ und der Bahn mit Wert s zu jedem Zeitpunkt t . Dieser Abstand hängt ab von der Differenz der Gravitationsfelder auf den beiden Bahnen. Die zweite Ableitung des Abstandes ist die Beschleunigung, mit der ein Beobachter, der auf $s = 0$ sitzt, das Teilchen s auf sich zu kommen sieht.

Die beiden Teilchenbahnen können weit auseinander liegen. Für sehr nahe verlaufende Nachbarbahnen gilt die Gleichung bis zur 1. Ordnung in s approximativ. Wir erkennen,

dass das Verhalten solcher Bahnen durch das Verhalten des Vektors $Z^i(t, 0)$ entlang der zentralen Bahn bestimmt wird und dieses wiederum ergibt sich durch Differenzieren nach s bei $s = 0$. Offensichtlich ist $Z^i(t, s) = \partial x^i(t, s)/\partial s$ ein Vektorfeld auf der Fläche S , welches an jedem Punkt einer Bahn in Richtung auf die 'infinitesimal benachbarte' Bahn zeigt. Man nennt dieses Vektorfeld auch **Jacobi-Feld** oder **Verbindungsvektor**.

Jacobi-Feld
Verbindungsvektor

Eine wichtige Eigenschaft eines Verbindungsvektors ergibt sich aus der Betrachtung von

$$\frac{\partial Z^i}{\partial t} = \frac{\partial^2 x^i}{\partial s \partial t} = \frac{\partial^2 x^i}{\partial t \partial s} = \frac{\partial U^i}{\partial s} \quad (6.1.3)$$

Dies bedeutet, dass die zwei Ableitungen $\partial/\partial s$ und $\partial/\partial t$ vertauschen und weil $\partial x^i/\partial s = Z^i$ und $\partial x^i/\partial t = \dot{x}^i =: U^i$ ist, folgt, dass das Verbindungsvektorfeld Z und das Tangentialvektorfeld U miteinander kommutieren.

Die Gleichung (6.1.2) für den Verbindungsvektor ist linear

$$\ddot{Z}^i = - \sum_k \Phi_{ik} Z^k.$$

Sein Verhalten wird also bestimmt durch die Matrix $\Phi_{ik} = \partial_{ik}\phi$, die Hesse-Matrix des Gravitationspotentials. Diese Matrix ist ein Maß für die *Inhomogenität* des Gravitationsfelds und damit verantwortlich für die Relativbeschleunigung zwischen verschiedenen frei fallenden Teilchen.

Um zu sehen, wie sich frei fallende Teilchen in einer relativistischen Raum-Zeit relativ zueinander bewegen, betrachten wir nun einen Schwarm von Geodäten. Wieder zeigt Abb. 6.1 das wesentliche. Wir interpretieren die Linien nun als eine 1-parametrische Schar von Geodäten, die eine Fläche S erzeugen, auf der jedes Raum-Zeit Ereignis durch die Parameter t und s eindeutig gekennzeichnet werden kann. Die zugehörigen Vektorfelder

$$Z = \frac{\partial}{\partial s}, \quad U = \frac{\partial}{\partial t}$$

kommutieren. Es gilt also $[Z, U] = 0$. Die Kurven mit konstantem s sind Geodäten, ihr Tangentialvektor erfüllt also die Gleichung

$$U^a \nabla_a U^b = 0.$$

Außerdem gilt wegen der Kommutativität der Vektorfelder

$$U^a \nabla_a Z^b = Z^a \nabla_a U^b.$$

Schreiben wir diese Gleichung in einem Koordinatensystem x^i auf, dann bekommen wir das genaue Analogon zu (6.1.3). Um ein Analogon zu (6.1.2) zu bekommen, be-

stimmen wir die zweite kovariante Ableitung von Z^a entlang U^a :

$$\begin{aligned}
U^a \nabla_a (U^b \nabla_b Z^d) &= U^a \nabla_a (Z^b \nabla_b U^d) \\
&= (U^a \nabla_a Z^b) \nabla_b U^d + U^a Z^b \nabla_a \nabla_b U^d = (Z^a \nabla_a U^b) \nabla_b U^d + U^a Z^b \nabla_a \nabla_b U^d \\
&= Z^a \nabla_a (U^b \nabla_b U^d) - Z^a U^b \nabla_a \nabla_b U^d + U^a Z^b \nabla_a \nabla_b U^d \\
&= U^a Z^b (\nabla_a \nabla_b U^d - \nabla_b \nabla_a U^d) = (U^a U^c R_{abc}{}^d) Z^b.
\end{aligned}$$

In dieser Rechnung wurde die Geodätengleichung, die Kommutativität der beiden Vektorfelder und die Ricci-Identität benutzt.

Schreiben wir *formal* $\dot{Z}^b := U^a \nabla_a Z^b$, dann haben wir

$$\ddot{Z}^d = (U^a U^c R_{abc}{}^d) Z^b.$$

Der Verbindungsvektor zwischen Geodäten genügt also einer ganz ähnlichen Gleichung wie der Verbindungsvektor zwischen Teilchenbahnen. An die Stelle der Hesse-Matrix des Gravitationspotentials tritt hier der $(1, 1)$ -Tensor

$$\Phi_a{}^b = U^c U^d R_{cad}{}^b,$$

der aus dem Riemann-Tensor und dem Tangentialvektorfeld an die Geodätenschar konstruiert wird. Senken wir den oberen Index mit der Metrik, so tritt eine weitere Analogie zutage. Wegen der Paarsymmetrie des Riemann-Tensors ist nämlich $\Phi_{ab} = \Phi_{ba}$, also ist dieser Tensor ebenso wie die Hesse-Matrix symmetrisch. Der Riemann-Tensor charakterisiert also die Inhomogenität des Gravitationsfeldes in enger Analogie zur Hesse-Matrix des Gravitationspotentials der Newtonschen Theorie.

Das Gravitationspotential genügt der Poisson-Gleichung

$$\Delta\phi = 4\pi G\rho.$$

Das Potential wird also ‘erzeugt’ durch die Massendichte ρ . Nun ist

$$\Delta\phi = \sum_i \frac{\partial^2 \phi}{\partial x^i \partial x^i} = \sum_i \Phi_{ii}$$

die Kontraktion (Spur) der Hesse-Matrix. Treiben wir die Analogie zum Höhepunkt, dann erwarten wir also für die Gleichung, die das Gravitationspotential in der ART – die Metrik – festlegt, eine Gleichung in der die Spur des Tensors Φ_{ab} eine Rolle spielt und in der die Quelle durch die *Energie* (Äquivalenz von Masse und Energie!) der vorhandenen Materie gegeben ist. Wir können auch nicht erwarten, dass wir mit einer Gleichung auskommen werden. Da ja das Gravitationsfeld der ART zehn Potentiale besitzt, erwarten wir auch zehn Gleichungen.

Die Spur des $(1, 1)$ -Tensors Φ_a^b ist

$$\Phi_a^a = U^c U^d R_{cad}^a.$$

Wir erwarten also eine Gleichung der Form

$$U^c U^d R_{cad}^a \approx 4\pi G \rho$$

wobei hier $\approx$ als 'so etwas wie' zu lesen ist.

Nun ist die linke Seite abhängig von dem beliebig gewählten Schwarm von Geodäten. Nehmen wir an, solche Schwärme verhalten sich in jeder Raumrichtung gleich, dann ist U^a ein beliebiger zeitartiger Vektor, den wir dann auf der linken Seite weglassen können.

Da die linke Seite von U^a abhängig ist, sollte dies die rechte Seite vernünftigerweise auch sein und zwar mit der gleichen Abhängigkeit von U^a wie die linke Seite. Wir erwarten also nicht für jedes U^a das gleiche ρ , sondern vielmehr, dass sich die rechte Seite in der Form

$$\rho = T_{ab} U^a U^b$$

schreiben lässt und dass demnach die Gleichungen der Gravitationstheorie die Form

$$R_{ab} \approx 4\pi G T_{ab}$$

haben sollten, wobei wir hier noch den **Ricci-Tensor** $R_{ab} = R_{acb}^c$ als Spur des Riemann-Tensors eingeführt haben. Der Ricci-Tensor ist symmetrisch, wie sich aus der Paarsymmetrie des Riemann-Tensors ergibt. Er hat zehn unabhängige Komponenten, also erhalten wir zehn einzelne Gleichungen. Auf der rechten Seite steht ebenfalls ein $(0, 2)$ -Tensor, von dem wir auch annehmen können, dass er symmetrisch ist. Dieser Tensor hat mit der Energie der Materie zu tun. Man nennt ihn **Energie-Impulstensor**.

Ricci-Tensor

Energie-Impulstensor

6.2 Der Energie-Impulstensor

Jede Materieform nimmt ein Raum-Zeit-Volumen ein. 'Massenpunkt' oder 'Punktladung' sind Idealisierungen, die nicht wirklich in der Natur vorkommen. (In der ART stellt sich sogar heraus, dass Massenpunkte mit der Theorie gar nicht verträglich sind). Materie kann daher nur als ein Kontinuum beschrieben werden.

Es gibt viele verschiedene Arten von Materie. Zwar setzt sich jede Materieverteilung letzten Endes aus Elementarteilchen zusammen, so dass es genügen würde, sich auf deren Beschreibung zu beschränken. Jedoch ist es hoffnungslos, die Eigenschaften einer Flüssigkeit aus denen der sie konstituierenden Elementarteilchen ableiten zu wollen. Ausserdem müsste man dann ein gültiges Modell für alle Elementarteilchen besitzen, was es derzeit noch nicht gibt. Daher werden *phänomenologische* Materiemodelle benutzt, mit denen man Flüssigkeiten, Gase, Plasmen usw. beschreibt.

Die Beschreibung dieser phänomenologischen Materiemodelle kann beliebig kompliziert werden, wie schon ein Blick auf die verschiedenen Gase zeigt, wo es mehrere Freiheitsgrade wie Rotations- und Vibrationsbewegung der einzelnen Moleküle geben kann. Festkörper können elastische Eigenschaften haben, d.h., sie können deformiert werden und auf diese Weise Energie und Impuls austauschen. All diesen Beschreibungen ist aber eines gemeinsam: sie gestatten uns, die Energie- und Impulseigenschaften der Materie zu quantifizieren. Es sind diese Eigenschaften einer Materieverteilung, die für uns hier relevant sind. In der SRT werden Energie und Impuls zu einem Vierer-Vektor, dem Energie-Impuls-Vektor, verknüpft, so dass die Energie beobachterabhängig wird. Mit einer Beschreibung des Energieinhalts einer Materieverteilung ist also immer auch eine Beschreibung des Impulsinhalts verbunden. Um Energie und Impuls voneinander zu trennen, muss ein Beobachter angegeben werden.

Da es sich bei einer Materieverteilung um ein ausgedehntes Kontinuum handelt, ist es sinnlos zu einem gegebenen einzelnen Ereignis die Energie oder den Impuls bestimmen zu wollen. Vielmehr ist es nur sinnvoll, von Energie- und Impuls*dichte* an einem Ereignis zu reden. Die gesamte Energie ergibt sich dann durch Integration über ein räumliches Volumen. Um also die Energie eines Systems zu bestimmen sind zwei Spezifikationen pro Ereignis notwendig, erstens ein Beobachter, also ein zeitartiger Einheitsvektor und zweitens ein räumliches Volumen über welches integriert werden soll. Dieses Volumen können wir uns an jedem Ereignis vorstellen als drei räumliche Einheitsvektoren, oder äquivalent dazu als ein zeitartiger Vektor, der auf den drei räumlichen Vektoren senkrecht steht. Wir kommen also zu dem Schluß, dass die Energiedichte an jedem Ereignis von zwei Vektoren abhängt und dass die Energie sich dann durch Integration der Energiedichte ergibt.

Diese Betrachtung macht plausibel, dass der Energieinhalt einer Materieverteilung durch einen $(0, 2)$ -Tensor T_{ab} , den Energie-Impulstensor, beschrieben wird. Ist u^a die 4-Geschwindigkeit eines Beobachters, dann ist

$$T_{ab}u^au^b$$

Energiedichte die **Energiedichte** der Materie, die der Beobachter in seinem Bezugssystem misst. Entsprechend ist der 4-Vektor

$$p_a = T_{ab}u^b$$

Energie-Impuls-Dichte die **Energie-Impuls-Dichte**, die der Beobachter der Materie zuweist. Jede Materieverteilung besitzt ihren eigenen Energie-Impulstensor, der sich aus den Materievariablen (Dichte, Druck, Geschwindigkeit, Ladung, Temperatur, Felder usw.) zusammensetzt. Allen Energie-Impulstensenoren sind aber zwei Eigenschaften gemeinsam.

Ein Energie-Impulstensor ist *symmetrisch*

$$T_{ab} = T_{ba}.$$

und *divergenzfrei*

$$\nabla_a T^{ab} = 0.$$

Die Divergenzfreiheit ergibt sich immer als Folge der Materie-Gleichungen, also der Gleichungen, denen die Materievariablen genügen.

Als Beispiel betrachten wir den Energie-Impulstensor von *Staub*, wechselwirkungsfreien Teilchen. Da sie nicht miteinander wechselwirken, bewegen sie sich auf Geodäten. Deren Tangentialvektor sei U^a . Der Energie-Impulstensor ist

$$T_{ab} = \rho U_a U_b.$$

Die Divergenz dieses Tensors ist der Vektor

$$\begin{aligned} D^b &= \nabla_a T^{ab} = (U^a \nabla_a \rho) U^b + \rho (U^a \nabla_a U^b) + \rho (\nabla_a U^a) U^b \\ &= (U^a \nabla_a \rho + \rho \nabla_a U^a) U^b = \nabla_a (\rho U^a) U^b. \end{aligned}$$

Um die Bedeutung von $D^b = 0$ zu sehen, schreiben wir $\nabla_a (\rho U^a) = 0$ für den Minkowski-Raum bzgl. kartesischer Koordinaten auf. Die 4-Geschwindigkeit u^a schreiben wir $\gamma(1, \mathbf{u})$ und für ∇_a setzen wir (∂_t, ∇) , wobei ∇ der dreidimensionale Ableitung $(\partial_x, \partial_y, \partial_z)$ ist. Damit ergibt sich

$$\partial_t (\rho \gamma) + \nabla (\rho \gamma \mathbf{u}) = 0.$$

Dies ist die **Kontinuitätsgleichung** für die Größe $\rho \gamma$, die Energiedichte des System für den Beobachter mit 4-Geschwindigkeit ∂_t . Dementsprechend ist $\nabla_a (\rho U^a) = 0$ die kovariante Form dieser Kontinuitätsgleichung. Ihre Bedeutung ist die *lokale Energieerhaltung*.

Kontinuitätsgleichung

Wir haben also folgendes Resultat: der Energie-Impulstensor für Staub ist divergenzfrei, denn die Staubteilchen erfüllen die lokale Energieerhaltung in Form der Kontinuitätsgleichung und die Weltlinien der Teilchen sind Geodäten.

6.3 Einstein-Tensor und Feldgleichung

Was wir aus dem vorigen Abschnitt 6.2 mitnehmen, sind die beiden Eigenschaften eines Energie-Impulstensors, seine *Symmetrie* und *Divergenzfreiheit*. Im Abschnitt 6.1 wurde die Form der Feldgleichung als

$$R_{ab} \approx T_{ab}$$

motiviert. Sowohl R_{ab} als auch T_{ab} sind symmetrische $(0, 2)$ -Tensoren. Jedoch ist R_{ab} im Gegensatz zu T_{ab} nicht divergenzfrei. Es stellt sich nun die Frage, ob es einen $(0, 2)$ -Tensor gibt, der allein aus Krümmungstermen konstruiert ist und der divergenzfrei ist. Zur Beantwortung dieser Frage kann uns nur die Bianchi-Identität helfen, denn sie ist die einzige Relation, die etwas über die Ableitung des Krümmungstensors aussagt. Sie lautet

$$0 = 3 \nabla_{[e} R_{ab]c}{}^d = \nabla_e R_{abc}{}^d + \nabla_a R_{bec}{}^d + \nabla_b R_{eac}{}^d.$$

Um eine Aussage für den Ricci-Tensor zu bekommen, kontrahieren wir diese Gleichung über die Indizes b und d und erhalten

$$0 = \nabla_e R_{ac} - \nabla_a R_{ec} + \nabla_b R_{eac}{}^b.$$

Eine weitere Kontraktion ergibt

$$0 = \nabla^c R_{ac} - \nabla_a R_c{}^c + \nabla^c R_{ac} = 2\nabla^b (R_{ab} - \frac{1}{2}g_{ab}R),$$

skalare Krümmung dabei ist $R = R_a{}^a$ die Spur des Ricci-Tensors, auch **skalare Krümmung** genannt. Damit haben wir einen Tensor gefunden, der unseren Kriterien – Symmetrie, Divergenzfreiheit und nur aus dem Krümmungstensor konstruiert – genügt. Dieser Tensor

$$G_{ab} = R_{ab} - \frac{1}{2}g_{ab}R$$

Einstein-Tensor heißt **Einstein-Tensor**. Er wurde von Einstein während seiner Suche nach den Feldgleichungen gefunden. Er postulierte daraufhin die Gleichung

$$G_{ab} = \kappa T_{ab},$$

wobei κ hier eine noch zu bestimmende dimensionsbehaftete Konstante ist. Offensichtlich folgt daraus

$$R_{ab} = \kappa (T_{ab} - \frac{1}{2}g_{ab}T),$$

wobei $T = T_a{}^a$ die Spur des Energie-Impulstensors ist.

Um die Konstante κ zu bestimmen, setzen wir den Energie-Impulstensor für Staub in diese Gleichung ein und überschieben mit der 4-Geschwindigkeit der Staubteilchen. Mit $T = \rho U_a U^a = \rho$ erhalten wir dann

$$\Phi_a{}^a = R_{ab}U^a U^b = \frac{\kappa}{2}\rho.$$

Nachdem Φ_{ab} in Analogie zur *negativen* Hesse-Matrix des Gravitationspotentials zu setzen ist, bekommen wir formale Übereinstimmung mit dem Newtonschen Gravitationsgesetz genau dann, wenn wir $\kappa = -8\pi G$ setzen².

Einsteinsche Feldgleichung Mit dieser formalen Betrachtung haben wir schließlich die **Einsteinsche Feldgleichung** in der Form

$$G_{ab} = -8\pi G T_{ab} \quad (6.3.1)$$

gefunden. Ein paar Bemerkungen hierzu

²Eine genauere Begründung beruht auf der Betrachtung des *Newtonschen Limes*, des Grenzfalles kleiner Geschwindigkeiten. In diesem Grenzfall dominiert die Zeit-Zeit-Komponente (R_{00}) des Ricci-Tensors alle anderen. Diese kann mit dem Gravitationspotential identifiziert werden. Sie erfüllt in dem genannten Grenzfall eine Poisson-Gleichung aus der sich dann der Wert für κ ablesen lässt. Dies ist genauer diskutiert in [12].

- Die angestellten Betrachtungen sind keine *Herleitung* der Gleichung. Dies ist nicht möglich. Es handelt sich vielmehr um eine Motivation auf Grund der erlangten Einsichten in das Verhalten der Gravitationswechselwirkung und unter Verwendung von *postulierten Prinzipien*, wie dem Äquivalenzprinzip.
- Die Einstein-Gleichung wird in dieser Form postuliert und muss sich experimentell bewähren.
- Es ist keineswegs so, wie oft vermutet, dass sich die ART mit dem Aufstellen der Einstein-Gleichung erschöpft hat und quasi abgeschlossen ist. Die Situation ist vielmehr so wie nach der Aufstellung der Maxwell-Gleichungen. Man interessiert sich jetzt für die *Lösungen* der Gleichung, um die *Phänomene* zu erforschen, die durch diese Gleichung beschrieben werden können. So wie in der Maxwell-Theorie verschiedene Ladungsmodelle (Leiter, Ferromagnet, Supraleiter, etc.) zu verschiedenen Phänomenen führen, so ist es das Ziel der Gravitationstheorie durch verschiedene Materiemodelle unterschiedliche Phänomene zu finden.
- Die Einstein-Gleichung bestimmt einen Teil der Raum-Zeit-Krümmung aus der *gegebenen* Materieverteilung. Sie sagt nichts aus über die Materieverteilung selbst. Dies muss zusätzlich spezifiziert werden, damit das System abgeschlossen ist. Dies bedeutet, dass man zusätzliche Gleichungen zu betrachten hat, die die Materievariablen (z.B. elektromagnetische Felder, Flüssigkeiten, elastische Körper o.ä.) bestimmen. Da diese Gleichungen üblicherweise Differentialgleichungen sind, in denen die kovariante Ableitung (und damit die Metrik, vgl. den Ausdruck für die Christoffel-Symbole) steht, erhält man auf diese Weise ein gekoppeltes System von Gleichungen, von denen die Einstein-Gleichung nur eine ist.
- Die Gleichung besteht (wie jede Gleichung) aus zwei Seiten, einer „himmlischen“ und einer „irdischen“ (Einstein). Die himmlische Seite ist die linke, krümmungsabhängige Seite. Es ist in der Tat eine erstaunliche Tatsache, dass es im Wesentlichen nur einen symmetrischen Tensor gibt, der divergenzfrei ist und aus höchstens zweiten Ableitungen der Metrik konstruiert werden kann. Die einzige mögliche Abänderung bestünde in der Addition eines Terms der Form Λg_{ab} mit der **kosmologischen Konstante** Λ . Bis heute ist nicht geklärt, ob die kosmologische Konstante positiv, negativ oder null ist³. Wir werden $\Lambda = 0$ annehmen.
- Die „irdische“ Seite enthält die Materieverteilung des betrachteten Systems und diese kann beliebig kompliziert werden. Um dieser Komplikation zu entgehen, werden wir im Weiteren nur noch die **Vakuum-Gleichung** betrachten, also die Gleichung

$$G_{ab} = 0.$$

³Es gibt jedoch neueste Ergebnisse aus der Vermessung der Inhomogenitäten der kosmischen Hintergrundstrahlung, die eine von Null verschiedene (negative) kosmologische Konstante nahelegen. Allerdings ist hier das letzte Wort noch nicht gesprochen.

7 Die Schwarzschild-Lösung

Nach dem Aufstellen der Feldgleichungen ist es uns jetzt natürlich ein Bedürfnis, Lösungen zu finden, um die neuen Phänomene, die durch die ART beschrieben werden, kennen zu lernen. Wir werden dies nicht im allgemeinsten Rahmen versuchen, denn dies würde zu einem hoffnungslos komplexen Unterfangen entarten¹. Vielmehr werden wir versuchen, durch vereinfachende Annahmen die Rechnung so einfach wie möglich zu machen.

Unser Ziel ist es zunächst, das Außenfeld einer kugelsymmetrischen und zeit-unabhängigen Quelle zu bestimmen. Wir werden uns also auf die Vakuum-Gleichungen beschränken. Was heißt es nun, eine 'Lösung der Feldgleichungen' bestimmen zu wollen. Eine Lösung der Feldgleichungen ist eine Raum-Zeit, besteht also aus einer Mannigfaltigkeit $\mathcal{M}$ und einer darauf definierten Metrik g_{ab} . Wir wissen, dass diese beiden koordinatenunabhängig definiert sind und dass es irrelevant ist für die physikalischen Konsequenzen, welche Koordinaten verwendet werden.

Es ist klar, dass wir Koordinaten verwenden müssen, wenn wir die Metrik aus den Feldgleichungen bestimmen wollen. Es gibt aber zwei Punkte, die wir dabei beachten müssen. Zwar ist die Wahl der Koordinaten irrelevant für die physikalischen Aussagen, sie ist jedoch sehr relevant für die Rechnungen, die durchzuführen sind. Es ist also von Vorteil, die Koordinaten so zu wählen, dass die Rechnung möglichst einfach wird. Zum anderen darf man nicht vergessen, dass die berechnete Metrik durch ihre Komponenten in dem verwendeten Koordinatensystem gegeben ist. Naturgemäß ist sie dann nur in der entsprechenden Koordinatenumgebung gekannt. Durch Lösen der Feldgleichungen bekommen wir also nicht *die* Lösung, sondern nur einen Teil davon, nämlich den Teil, der sich durch das gewählte Koordinatensystem beschreiben lässt. Nach dem Lösen der Gleichungen muss man also untersuchen, inwieweit die erhaltene Lösung 'vollständig' ist.

7.1 Herleitung der Schwarzschild-Metrik

Welche Konsequenzen hat nun unsere Beschränkung auf zeitunabhängige, kugelsymmetrische Felder? Beginnen wir mit der Zeitunabhängigkeit. Dies bedeutet, dass es eine

¹Die *allgemeine Lösung* der Einsteingleichung ist nach wie vor unbekannt.

Koordinate $x^0 = t$ gibt, so dass

$$\frac{\partial g_{ij}}{\partial t} = 0, \quad \text{für alle } i, j = 0 : 3.$$

Außerdem folgt, dass die Hyperflächen der Gleichzeitigkeit $t = \text{const}$ *raumartig* sind.

Fordern wir weiter, dass die Situation zeitumkehrinvariant ist, d.h. dass die Transformation $t \mapsto -t$ die Metrik invariant lässt, dann folgt, dass die Metrik die Form

$$g = g_{00} dt^2 - \sum_{i,j=1}^3 g_{ij} dx^i dx^j$$

haben muss. Insbesondere gibt es keine g_{0i} -Terme.

Die Kugelsymmetrie wird nun folgendermassen berücksichtigt. Wir betrachten eine beliebige Hyperfläche $\Sigma_t = \{t = \text{const.}\}$, also einen beliebigen Zeitpunkt. Wegen der Zeitunabhängigkeit der Metrik trifft die folgende Betrachtung dann für alle Zeitpunkte ebenso zu. ‘Kugelsymmetrie’ bedeutet, dass eine beliebige Drehung um ein festes Zentrum die Metrik invariant lässt. Wir nehmen ein beliebiges, aber festes Ereignis P in Σ_t heraus und lassen beliebige Drehungen auf P wirken. Dann durchlaufen die ‘gedrehten’ Ereignisse eine Kugelfläche. Machen wir diese Prozedur mit allen Ereignissen in Σ_t , dann kommen wir zu dem Ergebnis, dass jedes Ereignis in Σ_t auf genau einer Kugelfläche zu liegen kommt. Σ_t besteht also (wie eine Zwiebel) aus lauter ineinander verschachtelten Kugelflächen. Wir können also eine ‘radiale’ Koordinate R finden, derart dass $R = \text{const.}$ gerade eine dieser Kugelflächen ist.

Betrachten wir nun eine dieser Kugelflächen, die wir durch $t = \text{const.}$ und $R = \text{const.}$ bekommen. Diese ist eine ‘runde’ Kugel, also nicht nur eine 2-dimensionale Fläche mit der Topologie einer Kugel, sondern eine geometrische Kugel. Dies folgt aus der Tatsache, dass sie unter Drehungen in sich übergeht und dass die Metrik bei dieser Drehung sich nicht ändert. Nun wählen wir einen Punkt als Nordpol aus und führen die üblichen Polarkoordinaten ein, dann hat die Metrik auf dieser Kugel die Form

$$f(R)^2 (d\theta^2 + \sin^2 \theta d\phi^2)$$

mit einer unbekannten Funktion $f(R)$, welche die Eigenschaft besitzt, dass $4\pi f(R)^2$ gerade der Flächeninhalt der betrachteten Kugel ist. Diese Betrachtung gilt für jede der Kugelflächen, so dass die Raum-Zeit-Metrik die Gestalt

$$g = g_{00} dt^2 - g_{11} dR^2 - 2 g_{12} dR d\theta - 2 g_{13} dR d\phi - f(R)^2 (d\theta^2 + \sin^2 \theta d\phi^2)$$

bekommt. Dies ist noch nicht die einfachste Form. Wir haben nämlich noch nicht benutzt, dass die Wahl des Nordpols auf jeder Kugelfläche beliebig ist. Betrachten wir die Linie der Nordpole zu einem festen Zeitpunkt. Diese Linie ist also gegeben durch $t = \text{const}$ und $\theta = 0$ und ist immer transversal (nie tangential) zu den Kugelflächen.

Nun können wir aber durch eine geschickte Wahl der Pole auf jeder Kugel erreichen, dass diese Linie immer *senkrecht* zu den Flächen ist. Auf ähnliche Weise können wir erreichen, dass die Linie, die den Ursprung von ϕ auf dem Äquator entspricht (also durch $\phi = 0, \theta = \pi/2, t = \text{const.}$ gegeben ist) ebenfalls zu den Kugelflächen senkrecht ist. Dann haben wir erreicht, dass auf allen Kugelflächen 'das gleiche' Polarkoordinatensystem verwendet wird und die Metrik erhält ihre endgültige Gestalt

$$g = g_{00}(R) dt^2 - g_{11}(R) dR^2 - f(R)^2 (d\theta^2 + \sin^2 \theta d\phi^2).$$

Die metrischen Koeffizienten hängen nur noch von der radialen Koordinate R ab. Diese Koordinate ist noch nicht fixiert. Wählen wir eine andere radiale Koordinate r , so dass $R' = \partial R / \partial r \neq 0$ ist, dann ändert sich die Metrik in

$$g = g_{00}(R(r)) dt^2 - g_{11}(R(r)) R'(r)^2 dr^2 - f(R(r))^2 (d\theta^2 + \sin^2 \theta d\phi^2).$$

Sie behält also ihre Form bei, nur die Koeffizienten ändern ihre Abhängigkeit von den Koordinaten. Wir können diese Freiheit in der Wahl der Radialkoordinate ausnutzen, um $f(R(r)) = r$ zu setzen. Das bedeutet geometrisch, dass wir als Radialkoordinate $r = \sqrt{A/4\pi}$ setzen, wobei A der Flächeninhalt der Kugelfläche $r = \text{const}$ ist².

Üblicherweise schreibt man die metrischen Koeffizienten in etwas anderer Form, so dass sich schließlich die allgemeine statische und kugelsymmetrische Metrik ergibt

$$g = e^{2\nu} dt^2 - e^{2\lambda} dr^2 - r^2 (d\theta^2 + \sin^2 \theta d\phi^2). \quad (7.1.1)$$

Diese Metrik enthält nur noch zwei unbekannte Funktionen, ν und λ , abhängig von einer Variablen, r . Alle anderen Freiheitsgrade sind durch die Symmetrieannahme und die Wahl der Koordinaten eliminiert.

Offensichtlich gilt die oben angestellte Betrachtung auch für die flache Minkowski-Raum-Zeit, deren Metrik wir demnach als Spezialfall in der allgemeinen Form wiederfinden müssen. In der Tat ist dies für $\nu = \lambda = 0$ der Fall. Dann ergibt sich nämlich die Minkowski-Metrik in sphärischen Polarkoordinaten.

Mit dieser Metrik lässt sich nun der metrische Zusammenhang bestimmen. Die nicht verschwindenden Christoffel-Symbole sind

$$\begin{aligned} \Gamma_{tr}^t &= \nu', \\ \Gamma_{rr}^r &= \lambda', \quad \Gamma_{tt}^r = e^{2(\lambda-\nu)} \nu', \quad \Gamma_{\theta\theta}^r = -r e^{-2\lambda}, \quad \Gamma_{\phi\phi}^r = -r \sin^2 \theta e^{-2\lambda}, \\ \Gamma_{r\theta}^\theta &= \frac{1}{r}, \quad \Gamma_{\phi\phi}^\theta = -\cos \theta \sin \theta, \\ \Gamma_{\theta\phi}^\phi &= \cot \theta, \quad \Gamma_{r\phi}^\phi = \frac{1}{r}. \end{aligned}$$

²Wir hätten ebenso gut die Radialkoordinate so wählen können, dass $\sqrt{g_{11}(R(r))} R'(r) = 1$ wird.

und damit ergeben sich weiter (nach etwas länglicher Rechnung) die Komponenten des Ricci-Tensors

$$\begin{aligned}
 R_{tt} &= e^{2(\nu-\lambda)} \left(-\nu'' - (\nu')^2 + \nu'\lambda' - \frac{2}{r}\nu' \right), \\
 R_{rr} &= \left(\nu'' + (\nu')^2 - \nu'\lambda' - \frac{2}{r}\lambda' \right), \\
 R_{\theta\theta} &= e^{-2\lambda} \left(1 - e^{2\lambda} + r(\nu' - \lambda') \right), \\
 R_{\phi\phi} &= e^{-2\lambda} \left(1 - e^{2\lambda} + r(\nu' - \lambda') \right) \sin^2 \theta, \\
 R_{ij} &= 0 \quad \text{sonst.}
 \end{aligned} \tag{7.1.2}$$

Nun müsste man daraus noch den Einstein-Tensor G_{ab} berechnen, um die Feldgleichungen endgültig aufstellen zu können. Es gilt nun aber per definitionem

$$G_{ab} = R_{ab} - \frac{1}{2} R g_{ab} \iff R_{ab} = G_{ab} - \frac{1}{2} G g_{ab}, \tag{7.1.3}$$

so dass im Falle der Vakuumgleichungen

$$R_{ab} = 0$$

gilt. Die Feldgleichungen sind demnach den drei Gleichungen

$$\nu'' + (\nu')^2 - \nu'\lambda' + \frac{2}{r}\nu' = 0, \tag{7.1.4}$$

$$\nu'' + (\nu')^2 - \nu'\lambda' - \frac{2}{r}\lambda' = 0, \tag{7.1.5}$$

$$1 - e^{2\lambda} + r(\nu' - \lambda') = 0, \tag{7.1.6}$$

äquivalent. Dies sind drei Gleichungen für zwei Unbekannte. Das System scheint also überbestimmt zu sein. Dem ist aber nicht so. Dies folgt aus der Bianchi-Identität, mit deren Hilfe man zeigen kann, dass alle Gleichungen erfüllt sind, wenn man eine Lösung für zwei der Gleichungen gefunden hat.

Aus (7.1.4) und (7.1.5) folgt sofort $\nu' + \lambda' = 0$, also $\nu + \lambda = C$ für eine Konstante C . Setzen wir dies in (7.1.6) ein, ergibt sich die Gleichung

$$1 - e^{2\lambda} - 2r\lambda' = 0.$$

Multiplikation dieser Gleichung mit $e^{-2\lambda}$ zeigt die erstaunliche Konsequenz

$$\left[r(1 - e^{-2\lambda}) \right]' = 0,$$

also

$$e^{-2\lambda} = 1 - \frac{B}{r}$$

für eine weitere Konstante B und daraus folgt

$$e^{2\nu} = C^2 \left(1 - \frac{B}{r}\right),$$

wobei wir die Konstante C umdefiniert haben. Einsetzen in die drei Gleichungen zeigt, dass alle erfüllt sind, so dass wir die allgemeine Lösung gefunden haben. Die Metrik sieht also folgendermaßen aus

$$g = C^2 \left(1 - \frac{B}{r}\right) dt^2 - \frac{1}{1 - \frac{B}{r}} dr^2 - r^2 (d\theta^2 + \sin^2 \theta d\phi^2).$$

Wir haben nun nur noch die Konstanten zu interpretieren. Die Konstante C lässt sich einfach behandeln. Beachten wir, dass der erste Term der Metrik sich schreiben lässt als

$$C^2 \left(1 - \frac{B}{r}\right) dt^2 = \left(1 - \frac{B}{r}\right) d(Ct)^2,$$

so liegt es nahe, die Zeitkoordinate umzudefinieren, d.h., also die Koordinate $\bar{t} = Ct$ zu verwenden. Dies kommt effektiv der Wahl $C = 1$ gleich. Die Konstante C hat also keine physikalische Bedeutung sondern legt nur die Skala der Zeitkoordinate fest.

Im Gegensatz dazu ist die Konstante B von physikalischer Bedeutung. Es ist offensichtlich, dass $B = 0$ die Minkowski-Metrik zur Folge hat. $B \neq 0$ hat demnach etwas mit der Abweichung der Metrik von der flachen Metrik zu tun und diese kann nur von der Masse des felderzeugenden Objekts herrühren. Dass dies tatsächlich so ist, werden wir in Kürze sehen.

Wir schreiben daher $B = 2MG$ und setzen sofort die Gravitationskonstante $G = 1^3$. So erhalten wir die erste Lösung der Einstein-Gleichung, die berühmte **Schwarzschild-Metrik**

Schwarzschild-
Metrik

$$g = \left(1 - \frac{2M}{r}\right) dt^2 - \frac{1}{1 - \frac{2M}{r}} dr^2 - r^2 (d\theta^2 + \sin^2 \theta d\phi^2).$$

Offensichtlich hat diese Metrik Probleme: abgesehen von der Mehrdeutigkeit der Polarkoordinaten bei $\theta = 0, \pi$ gibt es zwei Stellen, an denen die Metrik nicht definiert ist, bei $r = 0$ und bei $r = 2M$. Das Problem bei $r = 0$ war zu erwarten, denn auch im Newtonschen Fall oder der Elektrostatik wird das Feld einer kugelsymmetrischen Quelle im Zentrum singular. Dies wird dann interpretiert, als ob im Zentrum eine 'punktförmige' Quelle (Massenpunkt oder Punktladung) säße, deren Dichte unendlich ist, da die Quelle keine Ausdehnung besitzt und deren Feld deshalb am Ort der Quelle divergiert. Die Singularität bei $r = 2M$ ist eine neue Qualität, die so nicht in der klassischen Theorie auftritt.

Schwarzschild-
Radius

Um ein Gefühl für diesen speziellen sog. **Schwarzschild-Radius** zu bekommen, be-

³Das heißt wir verwenden physikalische Einheiten in denen $c = 1$ und $G = 1$ gilt. Dieses sind die so genannten **geometrischen Einheiten**, in denen nicht nur Zeit und Länge in gleichen Einheiten gemessen werden, sondern ebenso auch die Masse.

rechnen wir seine Größe für einige realistische Massenwerte. Dazu muss man zuerst die physikalischen Konstanten G und c wiedereinführen. Dann lautet die Beziehung

$$\frac{R_s}{m} \approx 1,5 \cdot 10^{-27} \frac{M}{\text{kg}}.$$

Für die Sonne mit einer Masse von ca. $2 \cdot 10^{30}$ kg ergibt sich der Schwarzschild-Radius zu ca. 3 km, für die Erde hingegen ist der Schwarzschild-Radius 'nur' ca. 9 mm groß.

Die Metrik wird nicht nur für $M = 0$ flach, sondern auch für große Radien, also im Limes $r \rightarrow \infty$ ergibt sich die Minkowski-Metrik. Die Schwarzschild-Metrik ist demnach **asymptotisch flach**. Sie beschreibt ein System, dessen Feld mit zunehmendem Abstand abnimmt. asymptotisch flach

7.2 Geodäten in der Schwarzschild-Metrik

Wollen wir weitere Aussagen über die Phänomene in der Schwarzschild-Metrik bekommen, so müssen wir uns anschauen, wie sich die freien Teilchen in dieser Raumzeit verhalten, also die Geodätengleichung untersuchen.

Das Aufstellen dieser Gleichung ist einfach, wenn auch mühsam und ergibt das folgende System (Punkt bedeutet Ableitung nach dem affinen Parameter s)

$$\ddot{t} + 2\nu' \dot{t} = 0, \quad (7.2.1)$$

$$\ddot{r} + \lambda' \dot{r}^2 + e^{2(\nu-\lambda)} \nu' \dot{t}^2 - r e^{-2\lambda} (\dot{\theta}^2 + \sin^2 \theta \dot{\phi}^2) = 0, \quad (7.2.2)$$

$$\ddot{\theta} + \frac{2}{r} \dot{r} \dot{\theta} - \sin \theta \cos \theta \dot{\phi}^2 = 0, \quad (7.2.3)$$

$$\ddot{\phi} + \frac{2}{r} \dot{r} \dot{\phi} + 2 \cot \theta \dot{\phi} \dot{\theta} = 0. \quad (7.2.4)$$

Dies sind zunächst vier gewöhnliche Differentialgleichungen zweiter Ordnung, die sich aber reduzieren lassen. Zunächst bemerken wir, dass sich (7.2.1) als

$$\frac{d}{ds} (e^{2\nu} \dot{t}) = 0$$

und (7.2.4) als

$$\frac{d}{ds} (r^2 \sin^2 \theta \dot{\phi}) = 0$$

schreiben lassen. Außerdem folgt man auf Grund der Kugelsymmetrie, dass eine Geodäte, die in einem Punkt in einer bestimmten Richtung startet immer in der Ebene bleiben muss, die durch die Startrichtung und die radiale Richtung aufgespannt wird. Wir können aufgrund der Kugelsymmetrie annehmen, dass der Startpunkt der Geodäte und ihre Anfangsrichtung in der Äquatorebene liegen. Dann ist anfänglich $\theta = \pi/2$

und $\dot{\theta} = 0$. Nun ist aber $\theta = \pi/2$ eine Lösung von (7.2.3), die wegen der Eindeutigkeit die einzige Lösung zu den vorgegebenen Startwerten ist. Das heißt die Geodäte liegt immer in der Äquatorebene. Das bedeutet schließlich, dass wir (7.2.3) ignorieren dürfen, sofern wir $\theta = \pi/2$ setzen.

Ein letzter Punkt ergibt sich ganz allgemein aus der Geodätengleichung. Es gilt nämlich immer, dass das Betragsquadrat des Tangentialvektors v^a entlang einer Geodäte konstant ist. Dies folgt aus der kurzen Rechnung

$$v^a \nabla_a (v^b v_b) = (v^a \nabla_a v^b) v_b + v^b v^a \nabla_a v_b = 2(v^a \nabla_a v^b) v_b = 0,$$

wobei Verträglichkeit von Metrik und kovarianter Ableitung und die Geodätengleichung benutzt wurden.

Im vorliegenden Fall ist $v^a v_a = g_{ab} v^a v^b = e^{2\nu} \dot{t}^2 - e^{2\lambda} \dot{r}^2 - r^2 \dot{\phi}^2 = \text{const.}$, wobei wir schon $\theta = \pi/2$ verwendet haben. Damit haben wir das System von vier gewöhnlichen Differentialgleichungen 2. Ordnung auf drei Differentialgleichungen 1. Ordnung reduziert, die da lauten

$$\left(1 - \frac{2M}{r}\right) \dot{t} = E \quad (7.2.5)$$

$$r^2 \dot{\phi} = h, \quad (7.2.6)$$

$$\dot{r}^2 = E^2 - \left(1 - \frac{2M}{r}\right) \left(\frac{h^2}{r^2} + \varepsilon\right) \quad (7.2.7)$$

mit Konstanten E , h und ε . Letztere hat die Werte $1, 0, -1$ entsprechend der Wahl von zeit-, licht- und raumartigen Geodäten.

Wir interessieren uns zunächst für zeitartige Geodäten, also die Bahnen von frei fallenden Teilchen. Wir setzen daher $\varepsilon = 1$. Führen wir außerdem die Variable $u = 1/r$ ein so erhalten wir die Gleichung

$$\dot{u}^2 = u^4 \left((E^2 - 1) + 2Mu - h^2 u^2 + 2Mh^2 u^3 \right).$$

Sind wir nur an der Bahn des Teilchens interessiert, also an der Abhängigkeit $u(\phi)$ so ergibt sich für $\partial u / \partial \phi = \dot{u} / \dot{\phi}$ die Gleichung

$$\left(\frac{du}{d\phi} \right)^2 = \frac{1}{h^2} \left((E^2 - 1) + 2Mu - h^2 u^2 + 2Mh^2 u^3 \right) = F(u)/h^2. \quad (7.2.8)$$

Damit lässt sich die Bahn beliebig genau bestimmen. Leider ist es nicht möglich, die Gleichung in geschlossener Form zu lösen, weil die Integration auf elliptische Funktionen führt. Dennoch lassen sich qualitative Aussagen machen, wenn man das Polynom $F(u)$ genauer untersucht.

Wir sind am Verhalten der Bahn im Außenbereich interessiert, also für Werte $2M \leq r < \infty$ bzw. $0 < u \leq 1/(2M)$. Es ergibt sich

$$\begin{aligned} F(0) &= (E^2 - 1), & F(1/2M) &= E^2 > 0, \\ F'(0) &= 2M > 0, & F'(1/2M) &= 2M + \frac{h^2}{2M}, \\ F''(1/6M) &= 0. \end{aligned}$$

Damit gibt es vier qualitativ verschiedene Fälle:

1. $F(u)$ ist immer positiv. Dies entspricht einer Bahn, die vom Unendlichen ($u = 0$) kommend bis zu $r = 2M$ kommt.
2. $F(u)$ ist positiv für $0 < u < u_0$. Dies entspricht einer Bahn, die vom Unendlichen kommend bis zu einem nächsten Punkt gelangt, dort umkehrt und wieder ins Unendliche entweicht.
3. $F(u)$ ist positiv zwischen zwei Werten u_1 und u_2 . Dies ist eine gebundene Bahn, die zwischen zwei Extremalwerten hin und her oszilliert.
4. $F(u)$ ist nur in einem Bereich um $u = 1/2M$ positiv. Dies entspricht Bahnen, die vom Zentrum her kommend einen maximalen Abstand erreichen und dann wieder zurückfallen. Solche Bahnen sind auch in Fall 2 und 3 möglich.

Um diese Situation etwas quantitativer fassen zu können, wollen wir nun noch eine approximative Lösung der Gleichung (7.2.8) suchen. Dazu betrachten wir eine gebundene Bahn (also $u_1 < F(u) < u_2$) und nehmen an, dass der kubische Term klein gegenüber den anderen Termen sei. Nun ist

$$\frac{2Mh^2u^3}{h^2u^2} = \frac{2M}{r}$$

und

$$\frac{2Mh^2u^3}{2Mu} = \frac{h^2}{r^2} = r^2\dot{\phi}^2.$$

Die Approximation ist also gültig, wenn, erstens, die Bahn des Teilchens groß gegenüber dem Schwarzschild-Radius der Zentralmasse ist und wenn, zweitens, die Bahngeschwindigkeit klein gegenüber der Lichtgeschwindigkeit ist.

Dann haben wir die approximative Gleichung

$$\left(\frac{du}{d\phi}\right)^2 = \frac{1}{h^2} \left((E^2 - 1) + 2Mu - h^2u^2 \right) \quad (7.2.9)$$

zu lösen. Die Lösung mit der Nebenbedingung $du/d\phi = 0$ bei $\phi = 0$ ist durch die Kegelschnitt-Gleichung

$$u_0(\phi) = \frac{M}{h^2} (1 + e \cos \phi), \quad e^2 = 1 + (E^2 - 1) \frac{h^2}{M^2}$$

gegeben. Diese Gleichung beschreibt im vorliegenden Fall (zwei Nullstellen der rechten Seite von (7.2.9)) eine Ellipse.

Übung 7.1: Man zeige, dass für diese Ellipse mit Halbachsen a und $b \leq a$ die Relationen

$$b^2/a = \frac{h^2}{M} \quad e^2 = \frac{a^2 - b^2}{a^2}. \quad (7.2.10)$$

gelten.

Betrachten wir diese Lösung als die nullte Ordnung einer Störungsrechnung. Wir suchen jetzt die nächste Ordnung und setzen dazu die Lösung in der Form

$$u = u_0 + u_1$$

an. Setzen wir diesen Ansatz in (7.2.8) ein und behalten nur die linearen Terme in u_1 , dann erhalten wir

$$\sin \phi \frac{du_1}{d\phi} = \cos \phi u_1 - \frac{M^3}{eh^4} (1 + e \cos \phi)^3.$$

Diese Gleichung lässt sich leicht lösen und man erhält

$$u_1 = \frac{M^3}{h^4} \left[\frac{1 + 3e^2}{e} \cos \phi + 3 \left(1 + \frac{e^2}{2} \right) - \frac{e^2}{2} \cos 2\phi + 3e\phi \sin \phi \right].$$

Dies ist demnach die Abweichung der Bahn von einer Ellipse, die durch den kubischen Term hervorgerufen wird. Wir sehen hier verschiedene Terme, von denen der erste und der dritte periodisch und der zweite konstant ist. Ihr Effekt auf die Bahn ist ebenfalls periodisch, so dass die Bahn aufgrund dieser Störterme weiterhin geschlossen bleiben würde, gäbe es nicht den vierten Term. Dieser ist nicht periodisch und dies hat zur Folge, dass die Bahn nach einem Umlauf nicht an dem Punkt ankommt, von dem aus sie gestartet war, sie nicht mehr geschlossen. Diesen Effekt kann man berechnen, indem man die Umkehrpunkte der Bahn verfolgt, an denen demnach $du/d\phi = 0$ gilt. Für unsere approximative Lösung $u = u_0 + u_1$ ist

$$\frac{du}{d\phi} = -\frac{Me}{h^2} \sin \phi + \frac{M^3}{h^4} \left[-\frac{1}{e} \sin \phi + e^2 \sin 2\phi + 3e\phi \cos \phi \right].$$

Dieser Ausdruck verschwindet für $\phi = 0$. Die nächste Nullstelle ist nun aber nicht bei $\phi = \pi$, wie man nach der Newtonschen Theorie erwarten würde, sondern bei $\phi = \pi + \delta$, wobei δ durch Einsetzen in die Gleichung bestimmt werden kann. Wenn man nun die Näherungen $\sin(\pi + \delta) \approx -\delta$ und $\cos(\pi + \delta) \approx -1$ in die Gleichung einsetzt, ergibt sich

$$3\pi = \delta \left(\lambda^2 + \left(\frac{1}{e^2} + 2e + 3 \right) \right), \quad \text{mit } \lambda = \frac{h}{M}.$$

Wir nehmen nun an, dass die Halbachsen der ungestörten Bahn von ungefähr gleicher Größenordnung sind, also $b/a \approx 1$ und dass beide Halbachsen wesentlich größer als

der Schwarzschild-Radius des Zentralkörpers sind, also $a \gg M$, $b \gg M$ gilt. Dann gilt insbesondere $b/a \gg M/b$ und damit auch

$$\frac{h^2}{M} = \frac{b^2}{a} \gg M \Rightarrow \lambda^2 \gg 1.$$

Daher gilt für die Verschiebung des Umkehrpunkts nach einer halben Periode

$$\delta \approx 3\pi \frac{M^2}{h^2}.$$

Benutzen wir nun die Relationen (7.2.10), so finden wir, dass sich nach einer vollen Periode eine **Perihel-Drehung** von ca.

Perihel-Drehung

$$6\pi \frac{M}{a(1-e^2)} \doteq 6\pi \frac{GM}{a(1-e^2)c^2}$$

ergibt. Der zweite Ausdruck ist in physikalischen Einheiten aufgeschrieben. Einen Eindruck von der Größe dieser Verschiebung bekommen wir, wenn wir hier die Zahlen für Merkur (der innerste Planet hat die kleinste Halbachse) einsetzen. Die große Halbachse des Planeten beträgt 0.387 astronomische Einheiten⁴. Die Bahnexzentrizität ist $e = 0.205$ und die Umlaufzeit beträgt ca. 88 Tage. Mit diesen Zahlen ergibt sich eine Perihel-Drehung von ca. 43'' pro Jahrhundert.

7.3 Lichtstrahlen in der Schwarzschild-Metrik

Wir betrachten nun, wie sich Lichtstrahlen in der Schwarzschild-Metrik verhalten. Dazu müssen wir die Geodäten-Gleichung für lichtartige Geodäten untersuchen. Wir setzen also in (7.2.2) $\varepsilon = 0$ und gehen ansonsten ganz analog zum Falle der zeitartigen Geodäten vor. Dies führt uns auf die Bahngleichung

$$\left(\frac{du}{d\phi}\right)^2 = \left(\frac{E^2}{h^2} - u^2 + 2Mu^3\right) \quad (7.3.1)$$

für Photonen. Wieder können wir zunächst den kubischen Term als klein vernachlässigen (in diesem Fall heißt das, dass die Bahn weit weg vom Zentrum verläuft, so dass der Schwarzschild-Radius der Zentralmasse klein gegenüber dem Bahnradius ist) und erhalten so die approximative Lösung

$$u_0(\phi) = \frac{E}{h} \cos \phi$$

wobei wir die Integrationskonstante so gewählt haben, dass bei $\phi = 0$ der Umkehrpunkt der Bahn liegt. Diese Bahn liegt auf einer Geraden $r \cos \phi = L$ (vgl. Abb. 7.1), deren Abstand L vom Zentrum durch

$$L = 1/u_0(0) = h/E.$$

⁴1AE = $1.5 \cdot 10^8$ km

gegeben ist.

Nun suchen wir wieder die Gleichung für die kleine Störung, die durch den zusätzlichen kubischen Term verursacht wird. Wir setzen also wieder den Ansatz

$$u \approx u_0 + u_1 = \frac{E}{h} \cos \phi + u_1(\phi)$$

in die Gleichung ein und behalten nur die linearen Terme in u_1 . Dies führt auf die Gleichung

$$-\sin \phi u_1' = -\cos \phi u_1 + \frac{M}{L^2} \cos^3 \phi.$$

Dabei wurde der Term $3M/L^2 u_1$ ebenfalls vernachlässigt, weil wir annehmen, dass $2M \ll L$ ist, also die Bahn weit weg vom Schwarzschild-Radius verläuft.

Diese lineare Gleichung kann man lösen, indem man erst die homogene Gleichung löst und dann eine partikuläre Lösung bestimmt. Die homogene Gleichung lautet nach einer Umformung

$$\frac{u_1'}{u_1} = \frac{\cos \phi}{\sin \phi}$$

also

$$u_1(\phi) = C \sin \phi$$

mit einer beliebigen Konstanten C . Eine partikuläre Lösung kann man durch 'Variation der Konstanten' bestimmen, also durch den Ansatz

$$u_1(\phi) = C(\phi) \sin \phi.$$

Dies führt nach einer kurzen Rechnung auf die Gleichung

$$C' = -\frac{M \cos^3 \phi}{L^2 \sin^2 \phi} = -\frac{M}{L^2} \left(\frac{\cos \phi}{\sin^2 \phi} - \cos \phi \right),$$

deren Lösung sich durch Integration ergibt. Die allgemeine Lösung der gestörten Gleichung ist damit durch

$$u_1(\phi) = C \sin \phi + \frac{3M}{2L^2} \left(1 - \frac{1}{3} \cos 2\phi \right)$$

gegeben. Verlangen wir nun, dass die Bahn symmetrisch bzgl. $\phi = 0$ verlaufe, dann folgt $C = 0$ und die Bahngleichung der gestörten Bahn bekommt die Form

$$u = \frac{1}{L} \cos \phi + \frac{3M}{2L^2} \left(1 - \frac{1}{3} \cos 2\phi \right).$$

Die gestörte Bahn ist in Abb. 7.1 angedeutet.

Die Störung hat also eine Ablenkung der Bahn von einer Geraden zur Folge. Wir sprechen daher vom Effekt der **Lichtablenkung** im Gravitationsfeld. Zur Berechnung des

Lichtablenkung

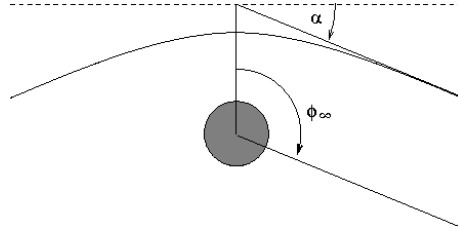


Abbildung 7.1: Zur Lichtablenkung

Ablenkungswinkels α benötigen wir den Grenzwinkel ϕ_∞ , der sich im Grenzfall $r \rightarrow \infty$ ($u = 0$) ergibt. Wir beachten, dass aufgrund der Geometrie $\phi_\infty = \alpha + \pi/2$ gilt. Setzen wir dies in die Bahngleichung ein, so folgt

$$0 = \cos(\pi/2 + \alpha) + \frac{3M}{2L} \left(1 - \frac{1}{3} \cos(\pi + 2\alpha) \right)$$

und für kleine Winkel α ergibt sich damit

$$\alpha \approx \frac{M}{L}$$

Die gesamte Ablenkung ist damit durch den doppelten Winkel gegeben. Setzen wir hier die Werte für die Sonne ein und wählen wir eine Bahn, die gerade noch die Sonnenoberfläche streift (also $L = R_\odot \approx 7 \cdot 10^5 \text{ km}$), dann ergibt sich ein Ablenkungswinkel von

$$2\alpha \approx 1,75''.$$

Dieser Wert wurde von Einstein berechnet und später von Eddington während einer totalen Sonnenfinsternis beobachtet. Heutige Messungen (Radioastronomie) bestätigen diesen Wert mit einer Genauigkeit von 1%.

Eine weitere einfache Konsequenz für Photonen ergibt sich, wenn wir ihre Frequenz bzw. Energie betrachten. Ein Photon bewegt sich mit Lichtgeschwindigkeit, seine Bahn ist ein lichtartige Geodäte mit lichtartigem Tangentialvektor k^a . Ein Beobachter mit Vierergeschwindigkeit u^a interpretiert das Skalarprodukt

$$g_{ab} k^a u^b$$

als die Energie e , bzw. mit der deBroglie-Beziehung

$$e = \hbar \omega$$

als Frequenz des Photons. Nehmen wir nun an, ein Photon bewege sich durch die Schwarzschild-Raumzeit und werde dabei durch ruhende Beobachter gemessen. 'Ruhend' soll hier bedeuten, dass die Vierergeschwindigkeit der Beobachter proportional

zum Killing-Vektor $\partial/\partial t$ sei, so dass für diese Beobachter nur die Zeit vergeht. Dann ist deren Vierergeschwindigkeit

$$u = \left(1 - \frac{2M}{r}\right)^{-1/2} \partial_t$$

und die Vierergeschwindigkeit des Photons hat die Form

$$k = \dot{t}\partial_t + \dot{r}\partial_r + \dot{\phi}\partial_\phi$$

so dass

$$\hbar\omega = g_{ab}u^a k^b = g(u, k) = \left(1 - \frac{2M}{r}\right)^{1/2} \dot{t} = \frac{E}{1 - \frac{2M}{r}}.$$

Also gilt

$$\omega \left(1 - \frac{2M}{r}\right)^{1/2} = E/\hbar = \text{const.}$$

Damit ergibt sich das Resultat, dass die Photonenenergie, die statische Beobachter messen, mit dem wachsendem Abstand vom Zentrum abnimmt, d.h. die Frequenz eines Photons, welches aus dem Gravitationsfeld herausläuft nimmt ab, es gibt eine **Rotverschiebung**.

Rotverschiebung

Zum Schluss betrachten wir noch radiale Photonen. Für diese gilt $\dot{\phi} = 0$ und wir erhalten die Gleichungen

$$\left(1 - \frac{2M}{r}\right) \dot{t} = E, \quad \dot{r}^2 = E^2.$$

Betrachten wir zuerst einlaufende Lichtstrahlen, so dass also $\dot{t} > 0$ und $\dot{r} < 0$ ist. Durch eine Umskalierung des affinen Parameters s können wir erreichen, dass $E = 1$ wird. Dann erhalten wir $\dot{r} = -1$ und die Lösung für *einlaufende* Lichtstrahlen

$$r = r_0 - s, \quad t = t_0 + s - 2M \log \frac{r_0 - s - 2M}{r_0 - 2M}.$$

Für *auslaufende* Lichtstrahlen haben wir $\dot{r} = 1$ und die Lösung

$$r = r_0 + s, \quad t = t_0 + s + 2M \log \frac{r_0 + s - 2M}{r_0 - 2M}.$$

Wir stellen also fest, dass die einlaufenden Geodäten nach einem endlichen Wert des affinen Parameters ($s = r_0 - 2M$) bei $r = 2M$ ankommen. Dort wird die entsprechende t -Koordinate jedoch unendlich. Es stellt sich jetzt aber die Frage, was passiert, wenn wir den affinen Parameter weiter erhöhen? Wo gehen die Geodäten hin?

Das gleiche Problem ergibt sich für die auslaufenden Geodäten, die alle vor einem endlichen Wert des affinen Parameters aus dem Bereich $r = 2M$ zu kommen scheinen. Dennoch divergiert auch hier die zugehörige t -Koordinate $t \rightarrow -\infty$. Wo waren die Geodäten vorher?

7.4 Ein schwarzes Loch

Um zu sehen, was bei $r = 2M$ passiert, müssen wir uns erst noch einmal vor Augen führen, wie wir zur Schwarzschild-Lösung gekommen sind. Die Symmetrie und Zeitunabhängigkeit haben uns auf eine bestimmte Form der Metrik geführt, deren Ricci-Tensor wir berechnet und die daraus folgenden Vakuum-Gleichungen wir gelöst haben. Dies liefert uns eine Raum-Zeit in einer bestimmten Koordinatenumgebung, hier die Koordinaten (t, r, θ, ϕ) . Es ist durch nichts garantiert, dass diese Koordinaten überall 'gut' sind. Offensichtlich ist die Metrik bei $r = 2M$ singulär. Ist dieses Verhalten aber auf die Raum-Zeit selbst zurückzuführen, oder liegt es an den Koordinaten?

Um diese Problematik etwas zu verdeutlichen betrachten wir ein sehr vereinfachtes Beispiel. Es sei die 1-dimensionale Mannigfaltigkeit $M = \mathbb{R}$ mit der Metrik

$$g = \frac{dy^2}{(1+y^2)^2}$$

gegeben. Dies ist für alle $y \in \mathbb{R}$ definiert. Betrachten wir nun die Geodätengleichung in diesem einfachen Fall, so folgt

$$\dot{y} = 1 + y^2,$$

also $y = \tan(s - s_0)$. Nun stellen wir auch hier fest, dass nach Ablauf eines endlichen Intervalls für den affinen Parameter die Koordinate y divergiert. Offensichtlich wollen die Geodäten 'über den Rand der Koordinatenumgebung hinaus'. Um zu sehen, was dann passiert, wählen wir den affinen Parameter s als Koordinate. Dann ergibt sich mit $y = \tan(s)$ und $-\pi/2 < s < \pi/2$ für die Metrik

$$g = ds^2.$$

Wir sehen also, die ursprüngliche Metrik bekommt in der neuen Koordinate eine andere Gestalt und die ursprüngliche Mannigfaltigkeit M erhalten wir zurück, wenn wir nur Werte von s im Intervall $]-\pi/2, \pi/2[$ betrachten. Dies muss aber nicht so sein, denn die 'neue' Metrik ist für beliebige Werte von s sinnvoll. Wir haben also die ursprüngliche Mannigfaltigkeit 'erweitert', indem wir zu einer geeigneten Koordinate übergegangen sind. Die neue Koordinate war der affine Parameter von Geodäten, welche wiederum selbst durch die Struktur der Metrik bestimmt sind.

Nehmen wir dieses Beispiel als Leitgedanken für die Untersuchung der Schwarzschild-Lösung in der Nähe von $r = 2M$. Betrachten wir einlaufende Lichtstrahlen. Wir hatten gesehen, dass für diese Strahlen

$$r = r_0 - s, \quad t + (r + 2M \log \frac{r - 2M}{r_0 - 2M}) = v = \text{const.}$$

Da entlang dieser Strahlen die t -Koordinate bei $r = 2M$ divergiert, während v konstant bleibt, wählen wir eine Koordinatentransformation, die t durch v und r ausdrückt

$$t = v - (r + 2M \log \frac{r - 2M}{r_0 - 2M}).$$

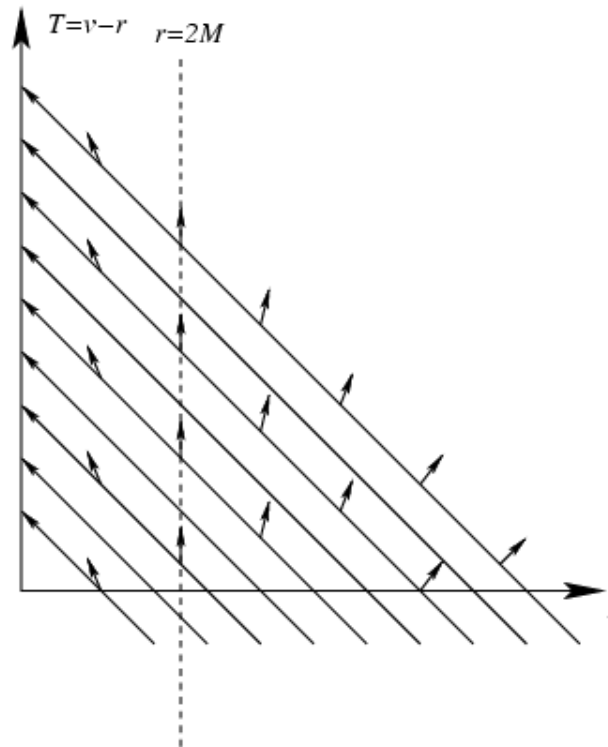


Abbildung 7.2: Die Raum-Zeit in Eddington-Finkelstein Koordinaten

Mit dieser neuen Koordinate wird die Metrik⁵

$$g = \left(1 - \frac{2M}{r}\right) dv^2 - 2dvdr - r^2 d\sigma^2.$$

Eddington-
Finkelstein
Koordinaten

Nun sehen wir, dass diese Metrik bei $r = 2M$ kein Problem mehr zeigt. Zwar verschwindet dort der Faktor vor dv^2 , aber das heißt nur, dass der entsprechende Koordinatenvektor ∂_v lichtartig wird. In diesem Koordinatensystem, den **einlaufenden Eddington-Finkelstein Koordinaten** sind die einlaufenden radialen Geodäten durch $v = \text{const.}$ charakterisiert. Jede einlaufende Geodäte erreicht irgendwann das Zentrum bei $r = 0$. Hier ist die Metrik nach wie vor singulär.

Auch hier haben wir die Mannigfaltigkeit 'erweitert', denn der Bereich, der von den (t, r) -Koordinaten überdeckt wird, reicht nur bis zu $r = 2M$, während die (v, r) -Koordinaten bis $r = 0$ reichen. Nun ist gezeigt, dass die Stelle $r = 2M$ ein vollständig regulärer Punkt der Raum-Zeit ist. Dennoch hat die Raum-Zeit dort eine spezielle Struktur. Wir betrachten Abb. 7.2, in der die Raum-Zeit in einem Diagramm dargestellt ist, in dem nach rechts die r -Koordinate aufgetragen wird und nach oben die 'Koordinate'

⁵Die Metrik der Einheitssphäre wird mit $d\sigma^2$ abgekürzt.

$T = v - r$. Dies hat zur Folge, dass die einlaufenden Geodäten alle auf Geraden unter einem Winkel von 45° laufen.

Außerdem ist gezeigt, wie sich die lokalen Null-Kegel verhalten. Diese erhält man dadurch, dass man an jedem Punkt der Raumzeit die beiden radialen lichtartigen Richtungen berechnet. Der Ansatz

$$l = a\partial_v + b\partial_r$$

führt auf die Gleichung

$$0 = g(l, l) = \left(1 - \frac{2M}{r}\right) a^2 - 2ab$$

mit den Lösungen $a = 0$ und $a/b = \frac{1}{2} \left(1 - \frac{2M}{r}\right)$. $a = 0$ entspricht der einlaufenden Richtung, denn es ändert sich nur r , während v (und die Winkel) konstant bleibt. Im Diagramm ist dies durch den Vektor $(1, -1)$, tangential an die einlaufenden Geraden, gegeben. Die auslaufende Richtung ist im Diagramm durch den Vektor

$$\left(1, \frac{r - 2M}{r + 2M}\right)$$

bestimmt. In großer Entfernung vom Zentrum verläuft die auslaufende Richtung ebenfalls unter 45° , jedoch nach auswärts in Richtung wachsende Radien. Je näher man zum Zentrum kommt, umso mehr 'klappt' die auslaufende Richtung nach innen. Dies ist natürlich der Effekt der Gravitation, die den Lichtstrahl 'bremst'. Dieser Effekt nimmt zu bis bei $r = 2M$ die r -Komponente verschwindet. Das bedeutet, dass die auslaufende Richtung nur eine Zeitkomponente besitzt und dies wiederum bedeutet, dass Lichtstrahlen, die bei $r = 2M$ radial nach außen laufen, tatsächlich nicht vom Fleck kommen, sondern immer bei $r = 2M$ bleiben. Geht man weiter zu Werten $r < 2M$, so wird die radiale Komponente der *auslaufenden* Richtung wie die der *einlaufenden* Richtung negativ. Das bedeutet, dass selbst Lichtstrahlen, die nach außen gerichtet sind, in Richtung kleiner Radien laufen. Sie können also dem Einfluss der Gravitation nicht entfliehen.

Die Fläche $r = 2M$ ist eine Trennfläche für das qualitativ verschiedene Verhalten der auslaufenden radialen Geodäten. Geodäten, die bei $r > 2M$ starten, kommen problemlos bis ins Unendliche. Hingegen sind Geodäten, die innerhalb starten, dazu verdammt, bis zum Zentrum bei $r = 0$ zu fallen. Da alle materiellen Teilchen auf zeitartigen Linien verlaufen, also sich an jedem Punkt der Raum-Zeit innerhalb des lokalen Null-Kegels fortbewegen, gilt dieses Verhalten auch für jede Materie. Man nennt die Fläche $r = 2M$ den **Ereignis-Horizont**, weil sie die Grenze der Raum-Zeit-Region ist, die für einen Beobachter im Unendlichen 'sichtbar ist'.

Ereignis-Horizont

Betrachten wir noch einmal ein radial auslaufendes Photon, welches bei einem Radius $R > 2M$ emittiert wird. Die Frequenz, die ein statischer Beobachter bei R misst sei ω_0 . Dieses Photon läuft ins Unendliche, wird dort ebenfalls von einem weiteren statischen Beobachter gemessen und hat eine Frequenz ω_∞ .

Nun hatten wir gesehen, dass

$$\omega \sqrt{1 - \frac{2M}{r}}$$

konstant ist. Damit gilt also

$$\omega_{\infty} = \omega_0 \sqrt{1 - \frac{2M}{R}}.$$

Das bedeutet aber, je näher das Photon bei $R = 2M$ emittiert wird, umso kleiner wird die Frequenz und damit die Energie des Photons, wenn es im Unendlichen ankommt. Im Grenzfall wird die Rotverschiebung unendlich. Das heißt aber, dass ein Beobachter der von außen ins Zentrum schaut, keine Information aus der Region $r \leq 2M$ bekommen kann. Sie ist für ihn 'schwarz'. Da diese Region wie eine Einweg-Membran agiert, durch die man hineinfallen kann, aber nicht mehr herauskommt, hat sich für diese Art von Phänomen der Name 'Schwarzes Loch' eingebürgert.

Was geschieht nun bei $r = 0$? Die Metrik ist singulär und um zu testen, ob es sich bei $r = 0$ um eine echte Singularität handelt, betrachtet man am einfachsten eine Größe, die unabhängig von jedem Koordinatensystem ist, und etwas über die Geometrie, also die Krümmung aussagt. Da der Ricci-Tensor identisch verschwindet (schließlich handelt es sich ja um eine Vakuum Raum-Zeit), bleibt nur der Riemann-Tensor selbst und als einfachste koordinatenunabhängige Größe bietet es sich an, den so genannten

Kretschmann-
Skalar

Kretschmann-Skalar, das Quadrat des Riemann-Tensors

$$R_{abcd}R^{abcd}$$

zu berechnen. Für die Schwarzschild-Lösung ergibt sich $48M^2/r^6$, also eine Funktion, die bei $r = 0$ divergiert. Damit ist geklärt, dass im Zentrum die Krümmung unbegrenzt anwächst. Damit wachsen die Gezeitenkräfte unbeschränkt, so dass jeder Körper, der zu nahe an die Singularität kommt, unweigerlich zerrissen wird.

Zum Schluss noch ein paar unsortierte Bemerkungen

- Wir hatten die Schwarzschild-Lösung mithilfe der einlaufenden Geodäten fortgesetzt und dabei eine Region II bekommen, die in der Zukunft der ursprünglichen Koordinatenumgebung, Region I, lag. Sozusagen die Region, in die die einlaufenden Geodäten *hinein*laufen. Man kann die gleiche Diskussion mit den auslaufenden Geodäten durchführen. Dann kann man die Mannigfaltigkeit um eine weitere Region III erweitern, nämlich diejenige, aus der die auslaufenden Geodäten herkommen. Diese ist nicht identisch mit der Region II, denn sie liegt in der Vergangenheit von Region I. Eine genauere Betrachtung der raumartigen Geodäten zeigt, dass es eine vierte Region geben muss, die in der Vergangenheit von Region II und in der Zukunft von Region III liegt. Diese so fortgesetzte Raum-Zeit heißt **Kruskal-Erweiterung**.

Kruskal-
Erweiterung

- Die Schwarzschild-Lösung ist eine kugelsymmetrische Lösung und es stellt sich die Frage, inwieweit das Auftreten einer Singularität und eines Ereignis-Horizonts

durch die Symmetrie bedingt ist. Es stellt sich heraus, dass unter sehr allgemeinen Bedingungen zwangsläufig eine Singularität entstehen muss, wenn der Energie-Inhalt einer genügend kleinen Region genügend hoch ist. Dann kollabiert diese Energie aufgrund ihrer selbstgravitierenden Wirkung und bildet eine Singularität.

- Die Schwarzschild-Lösung ist eine Lösung, bei der die Singularität hinter einem Ereignis-Horizont 'versteckt' ist. Dies muss nicht so sein. Wählt man beispielsweise in der Schwarzschild-Lösung den Massenparameter M negativ – aus verschiedenen Gründen eine unphysikalische Wahl – so bekommt man dennoch eine Lösung der Einsteingleichungen, die bei $r = 0$ singulär ist. Nun ist aber kein Horizont vorhanden. Aus verschiedenen Gründen hat Penrose die 'Hypothese der kosmischen Zensur' postuliert, wonach es keine 'nackten Singularitäten' geben soll, also Singularitäten, die für einen Beobachter im Unendlichen sichtbar sind. Singularitäten sollten dem gemäß immer hinter einem Ereignis-Horizont versteckt sein. Ob bzw. unter welchen Umständen dies so ist, ist eines großen ungelösten Probleme der klassischen ART.

8 Kugelsymmetrische statische Materieverteilungen

Nach der Behandlung der kugelsymmetrischen, statischen *Vakuum*-Gleichungen wenden wir uns jetzt den kugelsymmetrischen statischen Gleichungen mit rechter Seite zu. Lösungen dieser Gleichungen sind Kandidaten für Raumzeiten, die von statischen Sternen erzeugt werden.

8.1 Die grundlegenden Gleichungen

Wir haben die allgemeinen Feldgleichungen¹

$$G_{ab} = -8\pi T_{ab}$$

zu betrachten. Die erste Frage, die wir zu klären haben ist, wie wir die Materie, deren Energie-Impuls-Tensor auf der rechten Seite steht, beschreiben wollen. Weil eine wirklich exakte Beschreibung der Materie als eine Menge von Elementarteilchen nicht möglich und auch nicht nötig ist, wird im allgemeinen in der ART angenommen, dass die Materie sich durch eine Flüssigkeit und zwar eine **ideale Flüssigkeit** beschreiben lässt. Dies ist ein Schwarm von idealisierten (Punkt)teilchen, die eine 4-Geschwindigkeit u^a besitzen, und die miteinander wechselwirken. Ein Beispiel für eine ideale Flüssigkeit ist ein Gas, in dem z.B. van-der-Waals-Kräfte für die Wechselwirkung zwischen den einzelnen Molekülen sorgen. Diese Wechselwirkung führt zu einer Abhängigkeit zwischen Volumen und Druck des Gases, die im sog. pV-Diagramm dargestellt und durch eine **Zustandsgleichung** $p = p(V)$ beschrieben werden kann. Anstelle von Druck und Volumen verwendet man in der ART als Zustandsvariablen lieber den Druck p und die Energiedichte ρ des Gases, weil dies lokalisierte Größen sind, also an einem Raumzeitpunkt definiert sind.

Der Energie-Impuls-Tensor einer idealen Flüssigkeit ist durch den Ausdruck

$$T^{ab} = (\rho + p) u^a u^b - p g^{ab} \quad (8.1.1)$$

gegeben. Dabei bezeichnet ρ die Energiedichte, p den Druck und u^a die 4-Geschwindigkeit der Flüssigkeit. Die Energiedichte und der Druck sind durch eine Zustandsgleichung miteinander verknüpft.

¹Wir setzen der Einfachheit halber hier wieder $G = 1$.

Übung 8.1: Zeigen Sie, dass der $(1, 1)$ -Tensor T^a_b die Eigenwerte ρ und p besitzt. Mit welcher Vielfachheit kommen diese vor? Welchen kausalen Charakter haben die entsprechenden Eigenvektoren?

Übung 8.2: Zeigen Sie, dass aus der Divergenzfreiheit des Energie-Impuls-Tensors $\nabla_a T^{ab} = 0$ die **Euler-Gleichungen**

Euler-Gleichungen

$$\begin{aligned} u^a \nabla_a \rho + (\rho + p) \nabla_a u^a &= 0, \\ (\rho + p) u^b \nabla_b u^a + (g^{ab} - u^a u^b) \nabla_b p &= 0 \end{aligned} \quad (8.1.2)$$

folgen. Wie lassen sich diese Gleichungen interpretieren?

Tip: Zerlegen Sie $\nabla_a T^{ab} = 0$ in Anteile parallel und senkrecht zu u^a .

Wir betrachten nun die Feldgleichungen mit den Annahmen der Zeitunabhängigkeit und der Kugelsymmetrie. Die Symmetrieforderungen ergeben mit den gleichen Überlegungen wie im vorigen Kapitel die gleiche Form der Metrik, so wie sie in (7.1.1) angegeben wurde.

Für die rechte Seite der Gleichung, den Energie-Impuls-Tensor, ergeben sich ebenfalls einschränkende Bedingungen aufgrund der Symmetrie. Wir hatten gesehen, dass die Annahme der Zeitunabhängigkeit bedeutet, dass es eine Zeitkoordinate t gibt, so dass die Raumzeitmetrik von dieser Koordinate unabhängig ist. Betrachten wir einen Beobachter, der sich an einem festen Ort befindet, für den also nur die Zeit vergeht. Die Weltlinie $\gamma(\tau)$ dieses Beobachters wird dann beschrieben durch

$$\gamma(\tau) \doteq (t(\tau), x_0^1, x_0^2, x_0^3),$$

wobei x^1, x^2, x^3 räumliche Koordinaten sind. Aus dieser Darstellung folgt, dass die 4-Geschwindigkeit t^a des Beobachters proportional zum Koordinatenvektor ∂_t sein muss. Für einen solchen *statischen Beobachter* ändert sich die Raumzeit im Laufe der Zeit *nicht*. Insbesondere nimmt er auch keine Bewegung der Flüssigkeitsteilchen wahr. Die Geschwindigkeit v , die er der Flüssigkeit zuschreibt, ist durch die orthogonale Zerlegung

$$u^a = \gamma(t^a + v^a)$$

bestimmt. Dabei ist $\gamma = g_{ab} u^a t^b$, $g_{ab} t^a v^b = 0$ und $v^2 = -g_{ab} v^a v^b$ das Betragsquadrat der Geschwindigkeit. Da diese Geschwindigkeit verschwinden muss, ergibt sich, dass auch u^a proportional zum Koordinatenvektor ∂_t sein muss

$$u^a \doteq (u^0, 0, 0, 0).$$

Mit der kugelsymmetrischen Metrik (7.1.1) folgt aus der Normierungsbedingung

$$g_{ab} u^a u^b = 1 \iff e^{2\nu} (u^0)^2 = 1 \iff u^0 = e^{-\nu} \iff u_0 = e^{\nu}$$

Damit bekommt der Energie-Impuls-Tensor die Darstellung

$$T_{ab} \doteq \begin{pmatrix} \rho e^{2\nu} & 0 & 0 & 0 \\ 0 & p e^{2\lambda} & 0 & 0 \\ 0 & 0 & p r^2 & 0 \\ 0 & 0 & 0 & p r^2 \sin^2 \theta \end{pmatrix}$$

Übung 8.3: Bestimmen Sie die Komponenten der Gleichung $\nabla_a T^{ab} = 0$ für diesen Energie-Impuls-Tensor in einer statischen, kugelsymmetrischen Raumzeit. Zeigen Sie, dass die radiale Komponente als einzige nicht trivial ist und die Gleichung

$$p' + (\rho + p)v' = 0$$

liefert.

Mithilfe von (7.1.2) berechnen wir die skalare Krümmung

$$R = -2e^{-2\lambda} \left\{ v'' + (v')^2 - v'\lambda' + \frac{1}{r^2} (1 - e^{2\lambda}) + \frac{2}{r} (v' - \lambda') \right\} \quad (8.1.3)$$

und damit die Komponenten des Einstein-Tensors bzw. der Feldgleichungen

$$\begin{aligned} G_{tt} &= e^{2(v-\lambda)} \left\{ \frac{1}{r^2} (1 - e^{2\lambda}) - \frac{2}{r} \lambda' \right\} = -8\pi e^{2v} \rho, \\ G_{rr} &= -\frac{1}{r^2} (1 - e^{2\lambda}) - \frac{2}{r} v' = -8\pi e^{2\lambda} p, \\ G_{\theta\theta} &= -e^{-2\lambda} r^2 \left\{ v'' + (v')^2 - v'\lambda' - \frac{1}{r} (v' - \lambda') \right\} = -8\pi r^2 p, \\ G_{\phi\phi} &= \sin^2 \theta G_{\theta\theta} = -8\pi r^2 \sin^2 \theta p, \\ G_{ij} &= 0 \quad \text{sonst.} \end{aligned}$$

Ebenso wie im Vakuum-Fall so sind auch hier nur drei Gleichungen wesentlich. Sie lauten

$$\frac{1}{r^2} (1 - e^{2\lambda}) - \frac{2}{r} \lambda' = -8\pi e^{2\lambda} \rho, \quad (8.1.4)$$

$$\frac{1}{r^2} (1 - e^{2\lambda}) + \frac{2}{r} v' = 8\pi e^{2\lambda} p, \quad (8.1.5)$$

$$v'' + (v')^2 - v'\lambda' + \frac{1}{r} (v' - \lambda') = 8\pi e^{2\lambda} p. \quad (8.1.6)$$

Im Gegensatz zum Vakuum-Fall haben wir aber nun nicht nur zwei Funktionen zu bestimmen, sondern vier. Nun scheint das System unterbestimmt zu sein. Wir müssen aber berücksichtigen, dass p und ρ nicht unabhängig sind, sondern über die Zustandsgleichung miteinander verknüpft sind. Wir müssen also eine Zustandsgleichung wählen und können dann die drei Gleichungen benutzen um v , λ und ρ (oder p) zu bestimmen.

Bevor wir damit beginnen, ist es vorteilhaft, noch eine weitere Überlegung anzustellen. Wir wissen, dass die Divergenzfreiheit des Energie-Impuls-Tensors eine Konsequenz der Feldgleichungen ist. Das heißt, wenn wir die drei Gleichungen differenzieren und geeignet miteinander kombinieren, dann müssen wir die Komponenten der Gleichung

$$\nabla_a T^{ab} = 0$$

bekommen. Um eine aufwendige Rechnerei zu ersparen betrachten wir den folgenden Tensor

$$E^{ab} = \frac{1}{8\pi} G^{ab} + T^{ab}.$$

Es ist klar, dass E^{ab} verschwindet, wenn wir die Feldgleichungen erfüllen. Wir wollen dies aber im Moment noch nicht fordern. Dann gilt

$$\nabla_a T^{ab} = \nabla_a E^{ab},$$

weil der Einstein-Tensor immer per Konstruktion divergenzfrei ist. Wir nehmen einen beliebigen Kovektor v_a und berechnen

$$v_b (\nabla_a T^{ab}) = \nabla_a (v_b E^{ab}) - E^{ab} \nabla_a v_b.$$

Spezialisieren wir nun v_a als das Koordinatendifferential dr , dann wird aus dieser Gleichung

$$dr_b (\nabla_a T^{ab}) = \nabla_a (E^{ar}) - E^{ab} \nabla_a dr_b.$$

Mit den Christoffelsymbolen für die kugelsymmetrische, statische Metrik (7.1.1), die in Abschnitt 7.1 aufgestellt wurden, finden wir

$$\nabla dr = -\Gamma_{tt}^r dt^2 - \Gamma_{rr}^r dr^2 - \Gamma_{\theta\theta}^r d\theta^2 - \Gamma_{\phi\phi}^r d\phi^2,$$

so dass wir schließlich die Gleichung

$$dr_b (\nabla_a T^{ab}) = \nabla_a (E^{ar}) + \Gamma_{tt}^r E^{tt} + \Gamma_{rr}^r E^{rr} + \Gamma_{\theta\theta}^r E^{\theta\theta} + \Gamma_{\phi\phi}^r E^{\phi\phi}$$

Nun können wir folgendermaßen argumentieren: der erste Term auf der rechten Seite dieser Gleichung verschwindet wenn $E^{rr} = 0$ gilt, wenn also die Feldgleichung (8.1.5) erfüllt wird. Erfüllen wir außerdem die Feldgleichung (8.1.4), dann ist auch $E^{tt} = 0$ und, weil $E^{\phi\phi} = \sin^2 \theta E^{\theta\theta}$ gilt, haben wir die Aussage, dass $E^{\theta\theta} = 0$ gilt, genau dann wenn die r -Komponente der Divergenz des Energie-Impuls-Tensors verschwindet. Wir können also die Feldgleichung (8.1.6) durch die wesentlich einfachere Gleichung

$$p' + (\rho + p)v' = 0 \tag{8.1.7}$$

ersetzen.

8.2 Lösung der Gleichungen

Wir betrachten zuerst die Gleichung (8.1.4)

$$\frac{1}{r^2} (1 - e^{2\lambda}) - \frac{2}{r} \lambda' = -8\pi e^{2\lambda} \rho.$$

In dieser Gleichung kommt die Funktion v überhaupt nicht vor, sondern nur die Funktion λ . Die Gleichung wendet sich also nur an die Geometrie einer Fläche $t = \text{const}$ der Gleichzeitigkeit und sagt nichts aus über das Verhalten der Geometrie in zeitartigen Richtungen. Ein Umschreiben der Gleichung liefert

$$(1 - e^{-2\lambda}) + 2re^{-2\lambda}\lambda' = \left[r(1 - e^{-2\lambda}) \right]' = 8\pi r^2 \rho$$

und damit

$$e^{-2\lambda} = 1 - \frac{2m(r)}{r}, \quad (8.2.1)$$

wobei wir

$$m(r) = 4\pi \int_0^r s^2 \rho(s) ds + m_0$$

gesetzt haben. Die Konstante m_0 ist zunächst unbestimmt. Sie wird aber durch die folgende Überlegung festgelegt.

Wir betrachten eine Kugel S_r , gegeben durch alle Ereignisse mit konstanter Koordinate r in einer beliebigen Fläche konstanter Zeit t . Der Flächeninhalt dieser Kugel ist $A(r) = 4\pi r^2$; dadurch wurde ja die Koordinate r gerade festgelegt. Die Koordinate r ist nicht der geometrische Radius R der Kugel. Dieser ist gegeben durch den Abstand des Ursprungs $r = 0$ von einem beliebigen Punkt P auf der Kugel S_r , also durch die Länge der Kurve

$$\gamma(s) = (t, s, \theta_0, \phi_0), \quad s \in [0, r].$$

Übung 8.4: Zeigen Sie, dass diese Kurve eine Geodäte ist.

Dabei spielen die Werte der Winkel aufgrund der Kugelsymmetrie keine Rolle. Diese Länge ist gegeben durch das Integral

$$R(r) = \int_0^r \sqrt{g_{rr}(s)} ds = \int_0^r e^{\lambda(s)} ds.$$

Wenn wir nun diese Kugel schrumpfen, also $r \rightarrow 0$ durchführen, dann zieht sie sich auf den Punkt $r = 0$ zusammen. In einer regulären Raumzeit wird die Mannigfaltigkeit in der Umgebung eines Punktes beliebig genau durch den Tangentialraum approximiert. Die Raumzeitmetrik g_{ab} wird beliebig genau durch die flache Metrik des Minkowskiraums approximiert. Wir erwarten daher, dass sich die geometrischen Größen in der Umgebung eines solchen regulären Punktes immer mehr wie die entsprechenden Größe in einer flachen Raumzeit verhalten. Das bedeutet insbesondere, dass sich das Verhältnis von Oberfläche einer Kugel zum Quadrat ihres Radius im Grenzfall kleiner Radien immer mehr dem euklidischen Wert 4π annähern sollte. Daher erwarten wir, dass im Grenzfall

$$\lim_{r \rightarrow 0} \frac{4\pi R(r)^2}{A(r)} = 1$$

gilt. Wir berechnen diesen Grenzwert. Da Zähler und Nenner gemeinsam verschwinden, benutzen wir die Regel von l'Hôpital und erhalten

$$\lim_{r \rightarrow 0} \frac{4\pi R(r)^2}{A(r)} = \lim_{r \rightarrow 0} \frac{4\pi R(r)^2}{4\pi r^2} = \lim_{r \rightarrow 0} R'(0)^2 = e^{2\lambda(0)}.$$

Wir finden also, dass die Regularität der Raumzeit bei $r = 0$ erfordert, dass

$$\lambda(0) = 0 \iff m_0 = 0$$

gilt.

Eine weitere Eigenschaft der Funktion $m(r)$ bekommen wir, wenn wir uns vergegenwärtigen, dass die Flächen $t = \text{const}$ überall *raumartig* sind, und dies aufgrund der Statik der Raumzeit auch sein müssen.

Übung 8.5: Was geschieht, wenn die Flächen $t = \text{const}$ in einem Bereich nicht mehr raumartig sind? Warum widerspricht dies der Statik der Raumzeit?

Die Flächen sind genau dann raumartig, wenn ihr Normalenvektor zeitartig ist, also wenn der Vektor ∂_t zeitartig ist. Dies ist der Fall, wenn

$$g_{tt} = e^{2\nu} > 0. \quad (8.2.2)$$

Außerdem folgt aus der Raumartigkeit der Flächen, dass alle Tangentialvektoren an die Flächen raumartig sind, so dass also insbesondere der Vektor ∂_r raumartig ist. Dies ist dann der Fall, wenn

$$-g_{rr} = e^{2\lambda} = \left(1 - \frac{2m(r)}{r}\right)^{-1} > 0$$

ist. Wir bekommen daher die Beziehung

$$r > 2m(r),$$

die auf allen Flächen $t = \text{const}$ gelten muss.

Damit lautet die Funktion

$$m(r) = 4\pi \int_0^r s^2 \rho(s) ds = \int \rho(s) s^2 \sin \theta ds d\theta d\phi.$$

Dies sieht verdächtig nach einem Volumenintegral über ρ aus. Und weil ρ die Massen- (bzw. Energie-) dichte der Materie beschreibt, liegt die Bezeichnung **Massenfunktion** für $m(r)$ nahe. Eine genauere Betrachtung zeigt uns aber, dass das wirkliche Volumenintegral etwas anders aussieht, denn das geometrische Volumenelement auf einer Fläche $t = \text{const}$ ist durch

$$d^3V = e^{\lambda} r^2 \sin \theta dr d\theta d\phi$$

definiert.

Übung 8.6: Verifizieren Sie diesen Ausdruck.

Massenfunktion

Tatsächlich ist also die Masse des Materials, das innerhalb einer Kugel S_r enthalten ist, durch

$$M(r) = 4\pi \int_0^r e^{\lambda(s)} s^2 \rho(s) ds = 4\pi \int_0^r \frac{s^2}{\sqrt{1 - 2m(s)/s}} \rho(s) ds$$

gegeben. Für gewöhnliche Materie ist $\rho > 0$ und daher auch $m(r) > 0$. Daher ist

$$\sqrt{1 - \frac{2m(r)}{r}} < 1,$$

so dass die Beziehung

$$M(r) > m(r)$$

folgt.

Wenn wir nun annehmen, dass ρ ausserhalb einer Kugel S_R , also für $r > R$ verschwindet, dann muss die Raumzeit in diesem Aussengebiet durch eine kugelsymmetrische, statische Vakuumlösung, also eine Schwarzschildlösung mit einem Massenparameter m beschrieben werden. Die Stetigkeit der Metrik bei $r = R$ erfordert, dass $m = m(R)$ gelten muss (vgl. unten). Der Massenparameter m der Schwarzschildlösung ist aber genau die gesamte, in der Raumzeit enthaltene, Masse bzw. Energie. Wir bekommen also die Beziehung

$$m = m(R) < M(R).$$

Die Differenz zwischen $M(R)$, der Masse der innerhalb von S_R befindlichen Materie und m , der gesamten Masse in der Raumzeit, kann also nur dadurch zustande kommen, dass auch das Gravitationsfeld zur Gesamtenergie beiträgt. Da $m - M(R)$ negativ ist, handelt es sich beim gravitativen Beitrag um *gravitative Bindungsenergie*.

Wenden wir uns nun der zweiten Feldgleichung (8.1.5) zu. Unter Verwendung von (8.2.1) lautet sie

$$\nu' = \frac{4\pi r^3 + m(r)}{r(r - 2m(r))}. \quad (8.2.3)$$

Um zu sehen, was diese Gleichung physikalisch bedeutet, betrachten wir den Fall schwacher Gravitationsfelder. Dieser ist durch die Bedingungen

$$m(r) \ll r, \quad p \gg \rho$$

charakterisiert.

Übung 8.7: Was sagen diese Bedingungen physikalisch aus? Beachten Sie, dass ρ als Energiedichte auch die Ruhenergie enthält.

Unter diesen Bedingungen wird

$$\nu' \approx \frac{m(r)}{r^2}.$$

Dieses ist nichts anderes, als die (einmal integrierte) kugelsymmetrische Newtonsche Poisson-Gleichung

$$\Delta\varphi = \frac{1}{r^2} \left(r^2 \varphi' \right)' = 4\pi\rho.$$

Die Funktion ν ist also in diesem Fall in Analogie zum Newtonschen Gravitationspotential φ zu setzen. Dies ist aber nur in diesem, dem statisch, kugelsymmetrischen Fall so. Im Allgemeinen kann man aus der Metrik nicht ein einziges Potential extrahieren.

Ist die Energiedichte und damit der Druck bekannt, dann kann man (im Prinzip) die Gleichung (8.2.3) integrieren und das Potential ν bestimmen. Wie aber finden wir die Energiedichte?

Dazu betrachten wir nun die letzte Gleichung (8.1.7). Diese liefert

$$\nu' = -\frac{p'}{\rho + p},$$

und damit, eingesetzt in (8.2.3),

$$p' = -(\rho + p) \frac{4\pi r^3 + m(r)}{r(r - 2m(r))}. \quad (8.2.4)$$

Betrachten wir auch hier wieder den Fall schwacher Felder, so erhalten wir

$$p' = -\rho \frac{m(r)}{r^2}, \quad (8.2.5)$$

die Bedingung für hydrostatisches Gleichgewicht in der Newtonschen Theorie.

Übung 8.8: Betrachten Sie ein Flüssigkeitselement im infinitesimalen Volumen, das durch $[r, r + dr] \times [\theta, \theta + d\theta] \times [\phi, \phi + d\phi]$ definiert wird. Berechnen Sie im Rahmen der Newtonschen Theorie die Gravitationskraft und die hydrostatische Kraft, die auf dieses Element wirken. Berechnen Sie die entsprechenden Kräfte auf eine Kugelschale der Dicke dr . Zeigen Sie so, dass die Gleichgewichtsbedingung durch obige Gleichung gegeben ist.

Die Gleichung (8.2.4) heißt **TOV-Gleichung** (Tolman-Oppenheimer-Volkoff) für relativistisches hydrostatisches Gleichgewicht.

TOV-Gleichung

Wir vergleichen einen relativistischen mit einem newtonschen Stern, indem wir annehmen, dass bei einem festen Wert der r -Koordinate die Werte von ρ , p und m übereinstimmen. Dann stellen wir fest, dass der relativistische Stern, für den die TOV-Gleichung gilt, einen grösseren Druckgradienten besitzt. Der relativistische Ausdruck für den Druckgradienten ist grösser als der newtonsche Ausdruck, weil der Zähler grösser und der Nenner kleiner ist. Der Zähler ist deshalb grösser, weil in beiden Faktoren jeweils eine Korrektur durch den Druck hinzu kommt. Dies erklärt sich dadurch, dass die Energie der Wechselwirkung, die durch den Druck beschrieben ist, ebenfalls einer Masse äquivalent ist, die zur Gravitationskraft beiträgt. Dadurch erhöht sich der Druck im Inneren des Körpers und damit wiederum der Beitrag zur Gravitationskraft. Der Nenner im relativistischen Ausdruck ist aufgrund der geometrischen Korrektur jedoch kleiner als der Nenner im newtonschen Ausdruck. Dieser qualitative Unterschied zwischen diesen beiden Fällen trägt wesentlich dazu bei, dass es im relativistischen Fall bei gegebener Zustandsgleichung zu Grenzwerten für die Masse von Körpern kommt, während die Körpermasse im newtonschen Fall beliebig groß werden kann. Dies werden wir unten genauer diskutieren.

Fassen wir das bisher Gefundene zusammen. Die Gleichungen für eine innere Lösung der statischen, kugelsymmetrischen Feldgleichungen mit idealer Flüssigkeit mit der Zustandsgleichung $\rho = \rho(p)$ sind

$$m' = 4\pi r^2 \rho, \quad (8.2.6)$$

$$p' = -(\rho + p) \frac{4\pi p r^3 + m(r)}{r(r - 2m(r))}, \quad (8.2.7)$$

$$v' = \frac{4\pi p r^3 + m(r)}{r(r - 2m(r))}. \quad (8.2.8)$$

Dies ist ein System gewöhnlicher Differentialgleichungen, die mit den Anfangsbedingungen $p(0) = p_0$, $m(0) = 0$ und $v(0) = 0$ integriert werden. Um eine Lösung zu bekommen, müssen wir also einen Zentraldruck p_0 und eine Zustandsgleichung wählen, dann liefert uns das Gleichungssystem eine eindeutige Lösung. Die Metrik in diesem Teil der Raumzeit ist gegeben durch (7.1.1) mit λ aus (8.2.1).

Diese Gleichungen gelten nur so lange, bis der Rand R des Materiebereichs erreicht ist. Dies ist durch die Bedingung $p(R) = 0$ angezeigt. Für grössere Werte des Radius finden wir dann ein Vakuum und daher eine Schwarzschildlösung mit Massenparameter m . Die Übergangsbedingungen, die am Sternrand gelten müssen, folgen aus der Tatsache, dass sich die Geometrie dort stetig und differenzierbar verhalten soll.

Übung 8.9: Betrachten Sie Geodäten, die den Sternrand überqueren. Glattheit der Geometrie bedeutet (unter anderem), dass die Geodäten ohne Knick oder Sprung durch die Sternoberfläche durchlaufen. Leiten Sie aus diesen Bedingungen her, dass $m = m(R)$ gilt und dass für die radialen Ableitungen

$$g'_{rr}(R) = -\frac{\frac{2m}{R^2}}{1 - \frac{2m}{R}}, \quad g'_{tt}(R) = \frac{2m}{R^2}$$

gilt.

8.3 Diskussion der Lösungen

Wie wir oben schon diskutiert hatten, ist bei einem gegebenen Dichteprofil die rechte Seite der TOV-Gleichung betragsmässig immer grösser als die der newtonschen Gleichung (8.2.5). Das bedeutet, dass der Druckverlauf steiler ist, so dass der zentrale Druck im relativistischen Fall immer höher liegt als im newtonschen Fall. Das bedeutet, dass es im relativistischen Fall schwieriger ist ein Gleichgewicht zu erreichen als im newtonschen Fall.

Wir betrachten dies exemplarisch im Fall eines konstanten Dichteprofiles innerhalb des Sterns

$$\rho(r) = \begin{cases} \rho_0 & r \leq R \\ 0 & r > R \end{cases}.$$

Dann erhalten wir in beiden Fällen für die Massenfunktion

$$m(r) = \begin{cases} \frac{4\pi}{3}\rho_0 r^3 & r \leq R \\ \frac{4\pi}{3}\rho_0 R^3 & r > R \end{cases}$$

Die newtonsche Gleichgewichtsbedingung kann man leicht integrieren. Es ergibt sich unter der Verwendung der Bedingung $p(R) = 0$

$$p_N(r) = \frac{2\pi}{3}\rho_0^2(R^2 - r^2). \quad (8.3.1)$$

Die relativistische Gleichgewichtsbedingung kann man ebenfalls explizit integrieren. In diesem Fall bekommt man nach einer kurzen Umformung die Gleichung

$$-\frac{2}{p + \rho_0} + \frac{2}{p + \rho_0/3} = \left[\log \left(1 - \frac{8\pi}{3}\rho_0 r^2 \right) \right]',$$

aus der sich schließlich die Lösung mit der Randbedingung $p_E(R) = 0$ ergibt

$$p_E(r) = \rho_0 \frac{\sqrt{1 - \frac{2M}{R}} - \sqrt{1 - \frac{2Mr^2}{R^3}}}{\sqrt{1 - \frac{2Mr^2}{R^3}} - 3\sqrt{1 - \frac{2M}{R}}}. \quad (8.3.2)$$

Dabei wurde $M = 4\pi/3\rho_0 R^3$ gesetzt. Diese Lösung und die entsprechende Metrik wurde ebenfalls von K. Schwarzschild gefunden. Man nennt sie daher die *innere Schwarzschildlösung*.

Wir vergleichen den zentralen Druck in beiden Fällen. Im newtonschen Fall erhalten wir

$$p_N(0) = \frac{2\pi}{3}\rho_0^2 R^2,$$

während wir im relativistischen Fall den Ausdruck

$$p_E(0) = \rho_0 \frac{\sqrt{1 - \frac{2M}{R}} - 1}{1 - 3\sqrt{1 - \frac{2M}{R}}}$$

bekommen. Nun sehen wir den wesentlichen Unterschied in den beiden Fällen. Während der Zentraldruck im newtonschen Fall immer einen endlichen Wert besitzt, wird der Zentraldruck im relativistischen Fall unbegrenzt hoch, wenn der Nenner verschwindet, wenn also

$$1 - \frac{2M}{R} = \frac{1}{9} \iff \frac{M}{R} = \frac{4}{9}$$

wird. Das bedeutet demnach, dass *homogene Sterne mit Radius R und einer Masse $M \geq 4/9R$ im Rahmen der Einsteinschen Gravitationstheorie nicht existieren können*. Anders ausgedrückt: Die Masse eines homogenen Sterns mit der Dichte ρ_0 ist höchstens so groß wie

$$M_{\max} = \frac{4}{9} \frac{1}{\sqrt{3\pi\rho_0}},$$

Nun kann man einwenden, dass ein homogenes Dichteprofil nicht gerade physikalisch sinnvoll ist. Man kann aber zeigen, dass allein aus der Annahme, dass die Dichte im Stern nach aussen hin nicht zunimmt, folgt, dass die Masse des Sterns durch den obigen Wert begrenzt ist, und zwar unabhängig vom Druckverlauf.

Dass solche Massenschranken existieren müssen, folgt schon aus der Statik der Metrik. Denn die dafür notwendige Bedingung $g_{rr} < 0$, die besagt, dass der Koordinatenvektor ∂_r raumartig sein muss, liefert schon eine obere Massenschranke

$$0 < -g_{rr} = 1 - \frac{2M}{R} \iff M < R/2.$$

Die physikalische Erklärung dafür ist die, dass bei nicht Erfüllen dieser Bedingung die Masse so groß ist, dass der Gravitationskraft nicht mehr widerstanden werden kann.

Wir wollen nun das Argument für die Existenz einer solchen Massenschranke angeben. Dazu gehen wir zurück auf die Gleichungen (8.1.5) und (8.1.6). Deren Differenz hängt nicht vom Druckverlauf ab. Wir schreiben also

$$-\frac{1}{r^2} (e^{-2\lambda} - 1) - \frac{1}{r} e^{-2\lambda} \lambda' - \frac{1}{r} e^{-2\lambda} \nu' + e^{-2\lambda} \{ \nu'' + (\nu')^2 - \nu' \lambda' \} = 0.$$

Wir verwenden nun in den ersten beiden Termen den Ausdruck (8.2.1) für $e^{-2\lambda}$ und erhalten

$$-r \left(\frac{m(r)}{r^3} \right)' - \frac{1}{r} e^{-2\lambda} \nu' + e^{-2\lambda} \{ \nu'' + (\nu')^2 - \nu' \lambda' \} = 0.$$

Die beiden letzten Terme können nun zusammengefasst werden in der Form

$$-\frac{1}{r} e^{-2\lambda} \nu' + e^{-2\lambda} \{ \nu'' + (\nu')^2 - \nu' \lambda' \} = r e^{-(\nu+\lambda)} \left[\frac{\nu'}{r} e^{\nu-\lambda} \right]',$$

so dass wir schliesslich die Gleichung

$$\left(\frac{m(r)}{r^3} \right)' = e^{-(\nu+\lambda)} \left(\frac{\nu'}{r} e^{\nu-\lambda} \right)'$$

vorliegen haben.

Wir nehmen nun an, dass die Dichte im Stern nach aussen zu strikt abnimmt, also

$$\frac{d\rho}{dr} < 0.$$

Dann folgt, dass auch die mittlere Dichte, die proportional zu $m(r)/r^3$ ist, nach aussen

hin abnimmt², so dass $\left(\frac{m(r)}{r^3}\right)' < 0$ gilt. Damit folgt

$$\left(\frac{v'}{r}e^{\nu-\lambda}\right)' < 0.$$

Wir integrieren diese Ungleichung vom Sternrand $r = R$ ausgehend nach innen und bekommen damit die Ungleichung für $r < R$

$$\frac{(e^\nu)'(R)}{R}e^{-\lambda(R)} < \frac{(e^\nu)'(r)}{r}e^{-\lambda(r)}.$$

Nun benutzen wir die Übergangsbedingungen am Sternrand, die uns die radiale Ableitung von $e^\nu = \sqrt{g_{tt}}$ bei $r = R$ liefern. Damit können wir die linke Seite der Ungleichung berechnen und erhalten

$$\frac{m(R)}{R^3} < \frac{(e^\nu)'(r)}{r}e^{-\lambda(r)}.$$

Wir lösen nach $(e^\nu)'$ auf und integrieren noch einmal unter Verwendung von (8.2.1) über das Intervall $[0, R]$. So erhalten wir die Ungleichung

$$e^{\nu(R)} - e^{\nu(0)} > \frac{m(R)}{R^3} \int_0^R \left(1 - \frac{2m(r)}{r}\right)^{-\frac{1}{2}} r \, dr,$$

bzw.

$$e^{\nu(0)} < e^{\nu(R)} - \frac{m}{R^3} \int_0^R \left(1 - \frac{2m(r)}{r}\right)^{-\frac{1}{2}} r \, dr.$$

Um das Integral auf der rechten Seite abschätzen zu können, rufen wir uns nochmals die Annahme $\rho' < 0$ in Erinnerung. Aus dieser folgt nämlich, dass der Wert von $m(r)$ immer größer ist, als der Wert der sich für einen homogenen Stern gleicher Masse und gleichen Radius ergäbe, so dass also

$$m(r) > \frac{m(R)r^3}{R^3}$$

Übung 8.10: Zeigen Sie diese Aussage. Betrachten Sie dazu zwei Sterne mit gleicher Masse und gleichem Radius, von denen der eine homogen ist und für den anderen die Annahme $\rho' < 0$ gilt.

²Es gilt $(m(r)/r^3)' = 4\pi\rho(r)/r - 3m(r)/r^4 = 3/r^4(4\pi/3\rho(r)r^3 - m(r))$. Andererseits folgt aus $\rho'(r) < 0$, dass $\rho(r) < \rho(s)$ wenn $r < s$. Damit folgt

$$\frac{4\pi}{3}\rho(r)r^3 = 4\pi \int_0^r \rho(r)s^2 ds < 4\pi \int_0^r \rho(s)s^2 ds = m(r),$$

also $(m(r)/r^3)' < 0$.

Mit dieser Ungleichung wird das Integral

$$\int_0^R \left(1 - \frac{2m(r)}{r}\right)^{-\frac{1}{2}} r \, dr > \int_0^R \left(1 - \frac{2m(R)r^2}{R^3}\right)^{-\frac{1}{2}} r \, dr = \frac{R^3}{2m(R)} \left(1 - \sqrt{1 - \frac{2m(R)}{R}}\right).$$

Damit gilt nun

$$0 < e^{\nu(0)} < \sqrt{1 - \frac{2m(R)}{R}} - \frac{1}{2} \left(1 - \sqrt{1 - \frac{2m(R)}{R}}\right) = \frac{3}{2} \sqrt{1 - \frac{2m(R)}{R}} - \frac{1}{2}.$$

Damit sind wir am Ziel, denn es folgt nun wie im homogenen Fall die Relation

$$\frac{m(R)}{R} < \frac{4}{9}.$$

9 Kosmologische Lösungen der Feldgleichungen

Die Kosmologie, also die Lehre von der dynamischen Struktur unseres Universums, ist ein Wissenschaftszweig, der erst durch die Allgemeine Relativitätstheorie entstanden ist. Dies lag wohl hauptsächlich daran, dass in der vorrelativistischen Zeit die Meinung vorherrschte, das Universum, also die Arena, in der sich Sterne, Planeten usw. bewegen, selbst statisch und ohne Struktur sei. Erst die Untersuchung der Einsteinschen Feldgleichungen im kosmologischen Kontext brachte die Einsicht, dass dies nicht notwendigerweise der Fall sein muss.

Zwar gab es schon früh Einwände gegen diese Auffassung, sie konnten sich aber nicht durchsetzen, wahrscheinlich nicht zuletzt deshalb, weil es an Alternativen zur Newtonschen Theorie mangelte. Ein wesentlicher Einwand wurde von Chéseaux (1744) und pointierter von Olbers (1826) angeführt. Dieses **Olberssche Paradoxon** beruht auf der Beobachtung, dass der Nachthimmel dunkel ist. Olbers nahm an, dass das Universum ein unendlicher euklidischer Raum sei, statisch und schon immer existierte. Er ging weiterhin davon aus, dass die mittlere Zahl n von Sternen pro Einheitsvolumen und die mittlere Leuchtkraft L der einzelnen Sterne räumlich und zeitlich konstant sind, sofern die Mittelungen über genügend große Gebiete erfolgen. Betrachtet man nun eine Kugelschale mit Radius R und Dicke dR , so kann man die Leuchtkraft aller in dieser Schale enthaltenen Sterne berechnen und bekommt dafür den Ausdruck

Olberssche
Paradoxon

$$4\pi(nL)R^2dR.$$

Andererseits ist die Intensität aufgrund eines einzelnen Sterns, die das Zentrum der Schale erreicht, gegeben durch seine Leuchtkraft, aber abgeschwächt durch den Faktor $4\pi R^2$. Damit ist die Intensität aufgrund der Leuchtkraft der Kugelschale im Zentrum der Schale gegeben durch

$$I(R) = \frac{4\pi(nL)R^2dR}{4\pi R^2} = (nL) dR.$$

Addieren wir die Beiträge aller Kugelschalen um dieses Zentrum auf, dann bekommen wir eine unendlich große Intensität! Das bedeutet, dass wir zu jedem Zeitpunkt eine unendlich große Strahlung empfangen sollten. Dies ist offensichtlich nicht so.

Nun kann man einwenden, dass bei der obigen Argumentation vergessen wurde, die Absorption der Strahlung durch zwischenliegende Sterne zu berücksichtigen. Dies kann

aber getan werden und ändert die Schlußfolgerung insoweit, als nicht mehr eine unendliche Intensität, sondern ein endlicher Wert erreicht wird. Dieser ergibt sich jedoch zu ungefähr dem 40,000-fachen der Intensität des Sonnenlichts, wenn die Sonne im Zenith steht.

Olbers selber versuchte das Paradoxon durch die Annahme aufzulösen, dass interstellarer Staub die Strahlung absorbiert. Dies würde aber nur dazu führen, dass sich dieser Staub auf hohe Temperaturen aufheizen würde und schließlich ebenso strahlen würde wie die Sterne.

Eine mögliche Auflösung des Paradoxons besteht in der Aufgabe der Annahme, dass das Universum schon immer existiert habe. Zusammen mit der endlichen Lichtgeschwindigkeit führt dies nämlich dazu, dass die Strahlung von genügend weit entfernten Sternen (und das ist offensichtlich der überwiegende Anteil) bei uns noch gar nicht angekommen ist. In anderen Worten, das Universum hat es noch nicht geschafft, die unendliche Intensität herbeizuschaffen, die vom Olbersschen Argument gefordert wird.

Die gängige Auflösung des Paradoxon besteht jedoch darin, die Unveränderlichkeit des Universums aufzugeben und insbesondere die Expansion des Weltalls zu akzeptieren. Dies hat den Effekt, dass Lichtquelle und Empfänger relativ zueinander bewegt sind, so dass der Doppler-Effekt eine Verschiebung des Lichts in den langwelligen Bereich verursacht. In dieser Sichtweise hatte Olbers recht. Der Nachthimmel strahlt gleichförmig, jedoch hat die gleichförmige Strahlung, die **kosmische Hintergrundstrahlung** eine extreme niedrige Temperatur, gerade mal 2,73K. Nur der Expansion des Weltalls haben wir es also zu verdanken, dass wir nicht gekocht werden.

kosmische Hinter-
grundstrahlung

Es ist interessant zu sehen, dass die Untersuchung eine scheinbar unschuldigen Frage „Warum ist der Nachthimmel dunkel?“ darauf hindeutet, dass das Universum expandiert und auch die Existenz einer kosmischen Hintergrundstrahlung nahelegt.

9.1 Das kosmologische Prinzip

Der Startpunkt der Kosmologie beruht auf der Einsicht, dass Möglichkeiten direkter Beobachtung, die uns zur Verfügung stehen, verschwindend gering sind. Seit Beginn der Menschheit haben wir nur einen Einblick in eine verschwindend kleine raumzeitliche Umgebung unseres Universums erhalten. Zwar können wir mit unseren Teleskopen bemerkenswert weit in die Tiefen des Weltalls vordringen, trotzdem erhalten wir nur Informationen von einem kleinen Teil unseres Rückwärtslichtkegels, also der Menge aller Ereignisse, von denen aus Licht zu uns gesendet werden konnte. Dementsprechend beruht die Kosmologie weniger auf Beobachtung als zu einem wesentlich grösseren Teil auf „philosophischen Vorurteilen“.

Eine empirische Tatsache, die in die Grundlage der Kosmologie Eingang findet, ist die Beobachtung, dass uns das Universum auf großen Skalen **isotrop** erscheint, d.h., es

isotrop

sieht in jeder Richtung gleich aus. Wie kommen wir zu dieser Beobachtung? Die Daten über die Verteilung von Galaxien zeigen zwar, dass es auf verschiedenen Längenskalen zu Anhäufungen von Galaxien kommt, und dass es sogar riesige Gebiete gibt, in denen gar keine Galaxien vorkommen. Dennoch scheint die Galaxienverteilung auf den größten Skalen isotrop zu sein. Die Beobachtungen von Radioquellen zeigen eine Verteilung mit einer Isotropie von bis zu 5%. Ebenso sind die kosmische Röntgen- und γ -Strahlung, die aus verschiedenen Richtungen auf uns zukommen, isotrop unter 5%. Der überzeugendste Hinweis auf die Isotropie des Universums auf großen Skalen ist aber die Existenz des Mikrowellenhintergrunds, einer Strahlung mit dem Spektrum eines schwarzen Körpers mit einer Temperatur von 2,73 Kelvin. Diese Strahlung ist isotrop mit einer Genauigkeit von unter einem 1%.

Das philosophische Vorurteil, welches in der Kosmologie zur Anwendung kommt, ist eine Verallgemeinerung des Kopernikanischen Weltbilds. Damit soll ausgedrückt sein, dass wir nicht erwarten können, dass unsere Position im Universum eine herausragende Rolle spielen würde. Vielmehr ist es vernünftig anzunehmen, dass alle Punkte im Universum gleichberechtigt sind. Das heißt insbesondere, dass das Universum von jedem Punkt aus betrachtet gleich aussieht, also isotrop um jeden Punkt ist. Dies führt uns auf die Annahme, dass das Universum räumlich **homogen** sei.

homogen

Übung 9.1: Zeigen Sie, dass aus Isotropie um zwei Punkte A und B schon Homogenität folgt. Betrachten Sie dazu als Beispiel eine Funktion ρ im $\mathbb{R}^3$, die isotrop bzgl. A und B sei, d.h. der Wert von ρ an einem Punkt P hängt nur vom Abstand zwischen P und A bzw. B ab.

Tip: Betrachten Sie ρ auf einem beliebigen Kreis um A. Legen Sie einen Kreis um B, der diesen schneidet und variieren Sie dessen Radius. Was können Sie über die Werte von ρ im überstrichenen Gebiet sagen? Folgern Sie die Konstanz (also Homogenität) von ρ überall, indem Sie die Kreise beliebig variieren.

Wir werden im Folgenden das *kosmologische Prinzip* beherzigen, welches man wie folgt formulieren kann:

In jeder Epoche stellt sich das Universum von jedem Punkt aus im Mittel bis auf lokale Ungleichmäßigkeiten gleich dar.

Man beachte, dass sich Isotropie und Homogenität auf die räumlichen Aspekte beziehen.

Die Annahme von Homogenität und Isotropie dient zunächst einmal dazu, die resultierenden Gleichungen zu vereinfachen. Damit ist nicht gesagt, dass dieses Prinzip auf jeden Fall richtig sei. Vielmehr erlaubt es, die Metrik so weit zu vereinfachen, dass man die wenigen kosmologischen Beobachtungen (insbesondere über die Isotropie der kosmischen Hintergrundstrahlung) benutzen kann, um die Parameter, die noch in der Metrik stecken, zu bestimmen. Kompliziertere Metriken würden auf mehr Parameter führen, die man durch Beobachtungen noch nicht festlegen könnte.

9.2 Die Geometrie des Weltalls

Wie lassen sich die Annahmen von Homogenität und Isotropie mathematisch fassen? Zunächst bedeutet Homogenität so etwas wie: zu einem beliebigen gegebenen 'Zeitpunkt' sieht jeder 'Punkt im Raum' aus wie jeder andere. Wir stellen uns also vor, dass die Raumzeit, die unser Universum modelliert, in lauter raumartige Hyperflächen Σ_t unterteilt werden kann, von denen jede einen bestimmten Zeitpunkt t ('Epoche') darstellt. Diese Hyperflächen sind nun jeweils so beschaffen, dass alle ihre Punkte (die 'Punkte im Raum') gleichberechtigt sind. Dies bedeutet technisch gesprochen, dass man zu je zwei Punkten P und Q eine Abbildung finden kann, die P auf Q abbildet, und dabei aber die Metrik unverändert lässt ('sieht gleich aus') vgl. Abb. 9.1. Solche

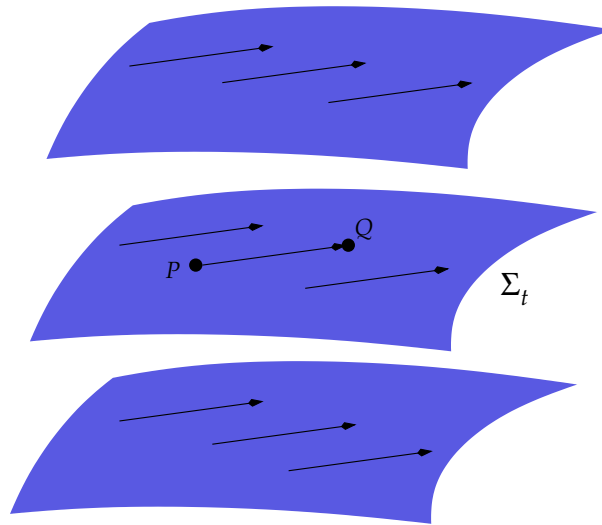


Abbildung 9.1: Die Flächen der Homogenität: zu jedem t und je zwei Punkten P und Q auf Σ_t gibt es eine isometrische Abbildung, die P auf Q abbildet.

Isometrie Abbildungen sind sogenannte **Isometrien**. Beispiele für eine Isometrie des flachen euklidischen Raums sind die Translationen und die Rotationen.

isotropen
Beobachter Kommen wir nun zur Isotropie. Zunächst folgern wir die Existenz der **isotropen Beobachter**, einer ausgezeichneten Schar von Beobachtern. Die isotropen Beobachter sind genau diejenigen, die das Universum tatsächlich als isotrop wahrnehmen. Denn nicht für jeden Beobachter ist dies tatsächlich der Fall: zwei Beobachter, die relativ zueinander in Bewegung sind, haben unterschiedliche Wahrnehmung (vgl. SRT). Wenn der eine Beobachter ein isotropes Universum sieht, so sieht der andere aufgrund der Aberration ein verzerrtes Bild des Universums. Ein anderes Argument beruht darauf, dass Materie, die sich im Universum befindet relativ zu einem isotropen Beobachter in Ruhe sein muss, während jeder andere Beobachter die Materie in Bewegung wahrnimmt, also eine bevorzugte Richtung feststellt.

Einen isotropen Beobachter beschreiben wir wie immer mit einer Weltlinie und ihrem Tangentialvektor u^a , der 4-Geschwindigkeit des jeweiligen isotropen Beobachters. Stellen wir uns vor, diese Beobachter sitzen an jedem Ereignis des Universums, dann liefert uns die jeweilige 4-Geschwindigkeit insgesamt ein ausgezeichnetes, zeitartiges Vektorfeld u^a , vgl. Abb. 9.2. Für einen isotropen Beobachter ändert sich nichts, wenn er sich

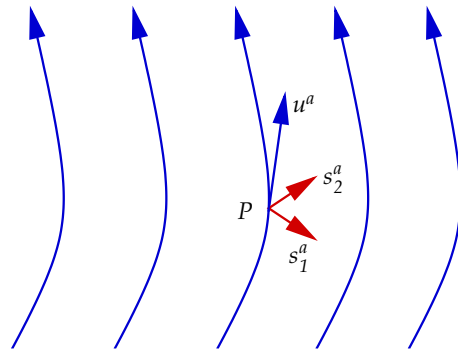


Abbildung 9.2: Weltlinien der isotropen Beobachter: zu jedem Ereignis P und je zwei räumlichen Richtungen s_1^a und s_2^a gibt es eine isometrische Abbildung, die P fest hält und s_1^a auf s_2^a dreht.

bei einem bestimmten Ereignis P von einer Richtung s_1 in eine beliebige andere Richtung s_2 dreht. Dies lässt sich wiederum mithilfe einer isometrischen Abbildung ausdrücken. Man sagt, die Raumzeit ist (räumlich) isotrop um das Ereignis P , wenn es eine isometrische Abbildung gibt, die den Punkt P auf sich abbildet, die 4-Geschwindigkeit u^a invariant lässt und einen Vektor s_1^a , der senkrecht auf u^a ist auf einen beliebigen anderen Vektor s_2^a dreht. Das bedeutet insbesondere, dass es keine geometrisch ausgezeichnete räumliche Richtung in P geben kann.

Daraus ergibt sich nun, dass die Weltlinien der isotropen Beobachter an jedem Punkt senkrecht auf den Flächen der Homogenität stehen. Denn wäre dies nicht der Fall, dann könnte man aus der nicht verschwindenden Projektion von u^a auf die Hyperflächen Σ_t eine bevorzugte Richtung konstruieren.

Übung 9.2: Bestimmen Sie den orthogonalen Projektor $P = P_a^b$ auf eine Hyperfläche Σ mit Normalenvektor n^a . Zeigen Sie explizit die Projektor-Eigenschaft $P^2 = P$. Bestimmen Sie ausserdem die induzierte Metrik h_{ab} , die von der Raumzeitmetrik g_{ab} auf Σ induziert wird. Welche Unterschiede ergeben sich, wenn Σ raum-zeit oder lichtartig ist?

Tip: Beachten Sie die Definition der induzierten Metrik $h(v, w) = g(Pv, Pw)$ für je zwei Vektoren v und w auf Σ .

Wir wollen nun die Form der Raumzeitmetrik g_{ab} bestimmen. Dazu führen wir geeignete Koordinaten ein. Wir wählen eine Hyperfläche Σ_0 und irgendwelche Koordinaten ξ^1, ξ^2, ξ^3 auf dieser Hyperfläche. Von jedem Punkt auf Σ_0 geht eine Weltlinie eines isotropen Beobachters aus. Auf jeder dieser Weltlinien wählen wir die Eigenzeit t des Beobachters als zeitliche Koordinate.

Betrachten wir eine beliebige Hyperfläche der Homogenität. Wir können jedem Punkt

auf Σ die Eigenzeit zuordnen die für den isotropen Beobachter vergangen ist, seit er auf Σ_0 gestartet war. Aufgrund der Homogenität von Σ (und der Tatsache dass es sich dabei um eine geometrische Konstruktion handelt) folgt, dass diese Funktion *konstant* sein muss. Die isotropen Beobachter stimmen also auf jeder Homogenitätshyperfläche hinsichtlich der Eigenzeitdauer überein. Das bedeutet, dass alle diese Hyperflächen durch $t = \text{const}$ gegeben sind.

Die räumlichen Koordinaten eines Ereignisses P legen wir wie folgt fest. Durch P bewegt sich ein isotroper Beobachter, d.h., P liegt auf der Weltlinie eines isotropen Beobachters, die wir zurückverfolgen, bis wir einen Punkt auf Σ_0 erreichen. Die Koordinaten dieses Punktes nehmen wir als räumliche Koordinaten von P .

Daraus ergibt sich, dass die Weltlinien der isotropen Beobachter durch die Gleichungen $\xi^1 = \text{const}$, $\xi^2 = \text{const}$ und $\xi^3 = \text{const}$ charakterisiert sind. Insbesondere ist die 4-Geschwindigkeit u^a gleich dem Koordinatenvektor $\partial/\partial t$.

Übung 9.3: Warum ist ∂_t ein Einheitsvektor?

Die Koordinatenlinien, die wir bekommen, indem wir t , ξ^2 und ξ^3 festhalten und nur ξ^1 variieren liegen offensichtlich in einer Homogenitätshyperfläche Σ_t . Das bedeutet, dass der Tangentialvektor an diese Linien, der Koordinatenvektor $\partial/\partial \xi^1$, tangential an Σ_t ist. Das gleiche folgt für die anderen beiden Koordinatenvektoren $\partial/\partial \xi^2$ und $\partial/\partial \xi^3$. Diese drei Vektoren bilden also in jedem Punkt eine Basis für den Tangentialraum an Σ_t .

Die Tatsache, dass u^a senkrecht auf den Σ_t steht, bedeutet, dass

$$g_{ab}u^a v^b = 0$$

für jeden Tangentialvektor v^b an Σ_t . Verwenden wir den Zusammenhang zwischen u^a und dem Koordinatenvektor ∂_t sowie die Tatsache, dass die räumlichen Koordinatenvektoren tangential an Σ_t sind, so ergibt sich

$$g_{0i} = g(\partial_t, \partial_{\xi_i}) = 0.$$

Damit lässt sich die Metrik nun in der Form

$$g = dt^2 - \sum_{i,k=1}^3 h_{ik}(t) d\xi^i d\xi^k \quad (9.2.1)$$

hinschreiben. Hier ist der räumliche Teil der Metrik die auf der Homogenitätshyperfläche Σ_t induzierte Metrik.

Unsere nächste Aufgabe besteht nun darin, die Auswirkung von Homogenität und Isotropie auf diese Metrik h_{ab} zu bestimmen.

9.3 Homogene und isotrope Mannigfaltigkeiten

Wir betrachten in diesem Abschnitt 3-dimensionale Mannigfaltigkeiten Σ mit einer Riemannschen Metrik h_{ab} . Wir beginnen mit der Feststellung, dass sich auch in diesem Fall der Riemannsche Krümmungstensor $R_{abc}{}^d$ sowie der Ricci-Tensor R_{ab} der Metrik einführen lässt (vgl. Abschnitt 5.4). Im 3-dimensionalen Fall gibt es jedoch eine Besonderheit: der Krümmungstensor ist durch den Ricci-Tensor eindeutig festgelegt. Mit anderen Worten, der Krümmungstensor enthält genauso viel Information wie der Ricci-Tensor. Tatsächlich besteht die Beziehung

$$R_{abcd} = R_{ac}h_{bd} - R_{ad}h_{bc} - R_{bc}h_{ad} + R_{bd}h_{ac} + \frac{1}{2}R(h_{ac}h_{bd} - h_{bc}h_{ad}) \quad (9.3.1)$$

zwischen den beiden Tensoren.

Übung 9.4: Zeigen Sie diese Beziehung, indem Sie wie folgt vorgehen.

1. Gehen Sie von der Identität (Beweis?)

$$\varepsilon_{abc}\varepsilon^{pqr} = 6\delta_{[a}^p\delta_b^q\delta_{c]}^r.$$

aus und zeigen Sie

$$\varepsilon_{abc}\varepsilon^{aqr} = \delta_b^q\delta_c^r - \delta_b^r\delta_c^q.$$

2. Definieren Sie den Tensor

$$S^{ab} = \varepsilon^{apq}R_{pqst}\varepsilon^{stb}$$

und zeigen Sie, dass dann

$$-4R_{ab} = S_{ab} - \frac{1}{2}h_{ab}S$$

gilt.

3. Drücken Sie R_{abcd} durch S_{ab} und damit durch R_{ab} aus.

Die Bedingung der Isotropie hatte zur Folge, dass es keine geometrisch ausgezeichneten Richtungen in der betrachteten Mannigfaltigkeit Σ geben darf. Betrachten wir nun den Ricci-Tensor als $(1,1)$ -Tensor $R^a{}_b$, also als lineare Abbildung im Tangentialraum $T_P\Sigma$ eines beliebigen Punkts auf Σ . Diese Abbildung ist symmetrisch, daher diagonalisierbar. Der Tangentialraum $T_P\Sigma$ zerfällt also in drei zueinander senkrechte Eigenräume zu den Eigenwerten dieser Ricci-Abbildung. Wären diese Eigenwerte nicht alle gleich, so gäbe es (mindestens) eine ausgezeichnete Richtung bei P , nämlich diejenige, die zum grössten (oder zum kleinsten) Eigenwert gehört. Dies ist im Widerspruch zur Isotropie. Folglich müssen die Eigenwerte alle gleich sein, das heißt die Ricci-Abbildung hat die Form $R^a{}_b = \lambda\delta_a^b$. Nehmen wir die Spur dieser Abbildung, so erhalten wir $3\lambda = R$, die skalare Krümmung. Und damit hat der Ricci-Tensor die Gestalt

$$R_{ab} = \frac{1}{3}Rh_{ab}.$$

Er ist proportional zur Metrik oder, wie man auch sagt, *reine Spur*. Daraus folgt für den Krümmungstensor

$$R_{abcd} = \frac{1}{6}R(h_{ac}h_{bd} - h_{bc}h_{ad}) \quad (9.3.2)$$

mit einer Funktion R , die zunächst noch vom jeweiligen Punkt $P \in \Sigma$ abhängen kann. Die Homogenität der Geometrie verlangt jetzt aber, dass R konstant sein muss. Dies folgt auch aus der Bianchi-Identität, denn sie liefert

$$0 = \nabla_{[e}R_{ab]}{}^{cd} = \frac{1}{3}\nabla_{[e}R\delta_a^c\delta_b^d]$$

woraus sich $\nabla_a R = 0$, also die Konstanz von R ergibt. Ein anderes Argument lässt sich wiederum mithilfe der Isotropie geben. Denn der Gradient $\nabla_a R$ würde wieder zu einer ausgezeichneten Richtung Anlass geben (nämlich diejenige der Vektoren, die senkrecht auf ihm stehen).

Damit haben wir also gefunden, dass die Metrik einer 3-dimensionalen, homogenen und isotropen Mannigfaltigkeit notwendigerweise einen Krümmungstensor mit der Form (9.3.2) mit einer Konstanten R besitzt. Solche geometrischen Räume nennt man auch *Räume mit konstanter Krümmung*. Wenn nun aber h_{ab} eine solche Metrik mit konstanter Krümmung ist, dann ist auch

$$\hat{h}_{ab} = \lambda^2 h_{ab}$$

Reskalierung für jedes konstante $\lambda \in \mathbb{R}$ eine solche Metrik. Es ist leicht zu sehen, dass diese **Reskalierung** der Metrik sich auf die Krümmungstensoren $\hat{R}_{abc}{}^d = R_{abc}{}^d$ nicht auswirkt. Daher folgt

$$\hat{R} = \lambda^{-2}R.$$

Der Effekt der Umskalierung der Metrik hat also zur Folge, dass sich die Konstante im Krümmungstensor mit einer positiven Konstante umskaliert. Dabei bleibt das Vorzeichen invariant. Es gibt also bis auf Reskalierung nur drei verschiedene Fälle zu untersuchen: $R > 0$, $R < 0$ und $R = 0$. Setzen wir $R = 6k$ mit $k = 0, \pm 1$, dann handelt es sich gerade um die Fälle

$$R_{abcd} = \pm(h_{ac}h_{bd} - h_{bc}h_{ad}), \quad R_{abcd} = 0.$$

Wir wollen in jedem der drei Fälle die entsprechende Metrik bestimmen. Dazu bemerken wir, dass eine Metrik, die um einen Punkt isotrop ist, zwangsläufig um diesen Punkt kugelsymmetrisch sein muss. Wir wählen uns also einen Punkt $O \in \Sigma$, erklären ihn zum Ursprung eines Polarkoordinatensystems und machen für die Metrik mit einer ähnlichen Begründung wie in Kap. 7 den Ansatz

$$h = A(r)^2 dr^2 + B(r)^2(d\theta^2 + \sin^2\theta d\phi^2).$$

Indem wir eine neue Radialkoordinate χ durch die Gleichung

$$\frac{d\chi}{dr} = A(r)$$

definieren, können wir diesen Ansatz noch etwas weiter vereinfachen und auf die Form

$$h = d\chi^2 + f(\chi)^2(d\theta^2 + \sin^2\theta d\phi^2)$$

bringen. Wir berechnen die skalare Krümmung dieser Metrik und erhalten nach etwas Rechnung

$$R = -2 \left(2 \frac{f''}{f} + \left(\frac{f'}{f} \right)^2 - \frac{1}{f^2} \right). \quad (9.3.3)$$

Um diese Gleichung zu lösen, betrachten wir zuerst den Ausdruck

$$\left(\frac{f'}{f^\alpha} \right)' = \frac{f''}{f^\alpha} - \alpha \frac{f'^2}{f^{\alpha+1}} = \frac{\alpha}{f^{\alpha-1}} \left(\frac{1}{\alpha} \frac{f''}{f} - \frac{f'^2}{f^2} \right).$$

Wir sehen, dass für $\alpha = -1/2$ die ersten zwei Terme in (9.3.3) bis auf einen gemeinsamen Faktor reproduziert werden. Setzen wir also (9.3.3) in diesen Ausdruck ein, so erhalten wir

$$-2f^{1/2} \left(f' f^{1/2} \right)' = \frac{R}{2} f^2 - 1.$$

Nach Multiplikation mit $f^{1/2}$ lässt sich diese Gleichung in der Form

$$- \left(f'^2 f \right)' = \frac{R}{2} f^2 f' - f'$$

schreiben und einmal integrieren.¹ Mit der Randbedingung $f(0) = 0$ verschwindet die freie Integrationskonstante und wir bekommen mit $R = 6k$

$$f'^2 = 1 - k f^2.$$

Diese Gleichung lässt sich jetzt für die drei Fälle gesondert lösen. Die Lösungen lauten (mit der Bedingung $f(0) = 0$)

$$f(\chi) = \begin{cases} \sin \chi & \text{für } k = 1, \\ \chi & \text{für } k = 0, \\ \sinh \chi & \text{für } k = -1. \end{cases}$$

Im Fall $k = 0$ lautet unsere Metrik daher

$$h_0 = d\chi^2 + \chi^2 (d\theta^2 + \sin^2\theta d\phi^2).$$

Es handelt sich also um die flache Metrik des $\mathbb{R}^3$ in Kugelkoordinaten. Der flache euklidische Raum ist somit ein Exemplar eines homogenen und isotropen Raumes.

¹Dabei setzen wir voraus, dass $f'(\chi) \neq 0$ ist. Wenn das nicht der Fall ist, dann ist $f = \text{const}$ und wir werden auf die sogenannte *Kantowski-Sachs Lösung* geführt. Diese ist aber physikalisch nicht sehr sinnvoll, so dass wir auf ihre weitere Diskussion verzichten.

Die anderen Fälle sind interessanter. Betrachten wir zuerst den Fall $k = 1$. Nun lautet die Metrik

$$h_+ = d\chi^2 + \sin^2 \chi (d\theta^2 + \sin^2 \theta d\phi^2).$$

Dies ist die Metrik der 3-dimensionalen Sphäre S^3 . Diese ist definiert als die Menge aller Einheitsvektoren im $\mathbb{R}^4$, also

$$S^3 = \{(u, x, y, z) \in \mathbb{R}^4 : u^2 + x^2 + y^2 + z^2 = 1\}.$$

Mit der normalen euklidischen Metrik im $\mathbb{R}^4$

$$g = du^2 + dx^2 + dy^2 + dz^2$$

induziert eine Metrik auf der S^3 . Wir führen Polarkoordinaten für die Punkte auf S^3 ein, indem wir setzen

dann ist $u^2 + x^2 + y^2 + z^2 = 1$, d.h., die so parametrisierten Punkte liegen auf der S^3 .

Übung 9.5: Zeigen Sie, dass die induzierte Metrik auf der S^3 in den angegebenen Koordinaten genau die Metrik h_+ ist.

Kommen wir schließlich zum Fall $k = -1$. Hier ist die Metrik

$$h_- = d\chi^2 + \sinh^2 \chi (d\theta^2 + \sin^2 \theta d\phi^2).$$

Offensichtlich gibt es eine enge Beziehung zwischen h_+ und h_- . Tatsächlich erhält man diese Metrik, indem man auf ganz analoge Weise wie bei $k = 1$ eine Untermannigfaltigkeit des $\mathbb{R}^4$ auswählt. Der einzige Unterschied besteht darin, dass nun nicht die euklidische Metrik g auf dem $\mathbb{R}^4$ betrachtet wird, sondern die Lorentz-Metrik

$$\eta = dt^2 - dx^2 - dy^2 - dz^2.$$

Man betrachtet nun das Einheitshyperboloid H^3 , welches durch

$$H^3 = \{(t, x, y, z) \in \mathbb{R}^4 : t^2 - x^2 - y^2 - z^2 = 1\}$$

definiert wird. Führt man nun wieder ganz analog zu oben Koordinaten ein, um das Hyperboloid zu parametrisieren, nämlich

$$x = \sinh \chi \sin \theta \cos \phi, \quad y = \sinh \chi \sin \theta \sin \phi, \quad z = \sinh \chi \cos \theta, \quad t = \cosh \chi$$

so ergibt sich wiederum $t^2 - x^2 - y^2 - z^2 = 1$ und die Lorentz-Metrik η induziert auf H^3 die eben berechnete Metrik h_- .

Übung 9.6: Zeigen Sie, dass man durch Einführen einer anderen Radialkoordinate $r = r(\chi)$, die Metriken $h_0, h_{\pm}$ auf die Form

$$\frac{dr^2}{1 - kr^2} + r^2(d\theta^2 + \sin^2 \theta d\phi^2)$$

bringen kann.

Die Annahme von Homogenität und Isotropie hat uns also zwangsläufig auf die obigen drei wesentlichen Geometrien geführt. Es handelt sich hierbei um die drei 3-dimensionalen Räume konstanter Krümmung, den euklidischen Raum, die 3-Sphäre S^3 und das raumartige Einheitshyperboloid H^3 . Dies sind bis auf Reskalierung mit einem positiven Faktor *die einzigen* 3-dimensionalen Mannigfaltigkeiten mit diesen Eigenschaften.

Was genau bedeutet nun die Tatsache, dass diese Räume homogen und isotrop sind? Wir hatten Homogenität und Isotropie mithilfe von isometrischen Abbildungen eingeführt. Wir wollen nun sehen, wie diese Abbildungen in den konkreten Beispielen wirken. Im Fall $k = 0$, dem flachen euklidischen Raum $\mathbb{E}^3$ ist das Wirken dieser Abbildungen am einfachsten zu sehen. Es ist offensichtlich, dass Translationen und Rotationen die Abstände von je zwei Punkten invariant lassen. Es handelt sich also dabei um isometrische Abbildungen. Umgekehrt kann man zeigen, dass jede isometrische Abbildung, die zusätzlich die Orientierung beibehält, sich aus einer Rotation und einer Translation zusammensetzt.

Übung 9.7: Zeigen Sie, dass es zu je zwei Punkten P und Q im $\mathbb{E}^3$ eine Isometrie gibt, die P auf Q abbildet. Zeigen Sie ausserdem, dass es zu je zwei Richtungen an einem Punkt P , repräsentiert durch je einen Einheitsvektor e_1 , e_2 eine Isometrie gibt, die P fest hält und e_1 auf e_2 dreht.

Im Fall $k = 1$, der Fall positiver Krümmung, betrachten wir zuerst den $\mathbb{R}^4$ mit seiner euklidischen Metrik. Auch dort sind Rotationen und Translationen die einzigen isometrischen Abbildungen. Die Rotationen um den Ursprung führen die S^4 in sich über. Wir können sie also als Abbildungen $S^4 \rightarrow S^4$ betrachten. Da sie die Metrik des $\mathbb{R}^4$ und damit die induzierte Metrik auf S^4 invariant lassen, sind sie isometrische Abbildungen der S^4 auf sich.

Im Fall $k = -1$, dem Fall negativer Krümmung, betrachten wir wieder zuerst den umgebenden Raum, den Minkowski-Raum. Auch hier gibt es isometrische Abbildungen. Diesmal handelt es sich um die Poincaré-Gruppe, die aus Lorentz-Transformationen und Translationen besteht.

Übung 9.8: Zeigen Sie, dass das Einheitshyperboloid H^3 unter der Lorentz-Gruppe invariant ist.

Die Invarianz des Einheitshyperboloids unter der Lorentz-Gruppe bedeutet wie im Fall positiver Krümmung, dass jede Lorentz-Transformation eine isometrische Abbildung des Einheitshyperboloids auf sich induziert.

Übung 9.9: Zeigen Sie, dass es zu je zwei Punkten P und Q in H^3 bzw. S^3 eine Isometrie gibt, die P in Q überführt. Bestimmen Sie die Transformationen, die P invariant lassen. Wie wirken diese auf die Tangentialvektoren an H^3 bzw. S^3 bei P ?

Tip: Wählen Sie Koordinaten so, dass P eine möglichst einfache Koordinatenbeschreibung erhält.

9.4 Kinematische Eigenschaften von kosmologischen Modellen

Nachdem wir die Auswirkungen der Forderung nach Isotropie und Homogenität des Universums untersucht haben, können wir die Metrik (9.2.1) nun noch genauer spezifizieren. Wir bekommen als Ansatz für die Metrik eines Universums

$$g = dt^2 - R(t)^2 \sum_{i,k=1}^3 h_{ik} d\xi^i d\xi^k = dt^2 - R(t)^2 \left(\frac{dr^2}{1 - kr^2} + r^2 d\theta^2 + r^2 \sin^2 \theta d\phi^2 \right) \quad (9.4.1)$$

wobei h_{ik} die metrischen Koeffizienten der Metriken h_0 oder $h_{\pm}$ sind. Der durch die Symmetrie nicht bestimmbare Skalenfaktor wird mit $R(t)$ bezeichnet. Er kann natürlich von einer Homogenitätshyperfläche zur anderen variieren. Diese Form der Metrik nennt man **Friedmann-Robertson-Walker-Metrik**.

Friedmann-
Robertson-Walker-
Metrik

Übung 9.10: Zeigen Sie, dass zwischen der 4-Geschwindigkeit u^a und der kosmologischen Zeit t die Beziehung

$$u_a = \nabla_a t$$

besteht. Wie lautet die Koordinatendarstellung von u_a ? Was können Sie über die kovariante Ableitung $\nabla_b u_a$ von u_a sagen?

Um die Bedeutung des Skalenfaktors R zu sehen, betrachten wir zwei isotrope Beobachter P und Q , die zu einer bestimmten Zeit $t = t_0$ durch die Homogenitätshyperfläche Σ_{t_0} in den Punkten p_0 bzw. q_0 hindurch laufen. Zu diesem Zeitpunkt beträgt ihr räumlicher Abstand $d(t_0) = R(t_0)s_0$, wobei s_0 der Abstand der beiden Durchtrittspunkte p_0 und q_0 ist, der mit der Metrik h auf der Homogenitätshyperfläche gemessen wird. Zu einem späteren Zeitpunkt, wenn die Beobachter P und Q die Homogenitätshyperfläche Σ_t durchlaufen, beträgt ihr Abstand $d(t) = R(t)s_0$, denn die räumlichen Koordinaten der Beobachter bleiben konstant, so dass auch der mit der räumlichen Metrik h gemessene Abstand unverändert bleibt. Die zeitliche Änderung des Abstands zwischen zwei Beobachtern kommt also nur durch die Veränderung des Skalenfaktors $R(t)$ zustande. Im Skalenfaktor kommt also die *relative Bewegung* der isotropen Beobachter zum Ausdruck.

Soweit zur Metrik. Kommen wir nun zur rechten Seite der Einstein-Gleichungen, dem Materie-Inhalt des Universums. Hier haben wir einen Energie-Impuls-Tensor anzugeben. Ebenso wie die Metrik, so ist auch dieser durch die Forderung nach Homogenität und Isotropie weitgehend eingeschränkt. Die einzigen 'Bausteine', die wir zur Verfügung haben, sind die Metrik g_{ab} , die 4-Geschwindigkeit u^a sowie Funktionen, die nur von der Zeit abhängen dürfen. Damit bleibt für den Energie-Impuls-Tensor nur die Gestalt $T_{ab} = A(t)u_a u_b + B(t)g_{ab}$ übrig. Wenn wir die Funktionen A und B 'umtaufen', erhalten wir die Form

$$T_{ab} = (\rho + p)u_a u_b - p g_{ab} \quad (9.4.2)$$

mit zeitabhängigen Funktionen ρ und p . Wir erhalten also zwingend die Form eines Energie-Impuls-Tensors für eine ideale Flüssigkeit. Die 4-Geschwindigkeit der Flüssig-

keit stimmt mit der 4-Geschwindigkeit der isotropen Beobachter überein, d.h., die isotropen Beobachter sind bezüglich der Flüssigkeit in Ruhe. Wir können die 'Flüssigkeit' als ein Gas interpretieren, das man erhält, wenn man die Verteilung der Galaxien und Galaxienhaufen im Universums ausschmiert und geeignet mittelt. Ein 'Flüssigkeitsteilchen' entspricht also einer einzelnen Galaxis.

Wie schon im Kap. 8 bekommen wir aus der Divergenzfreiheit des Energie-Impuls-Tensors

$$\nabla_a T^{ab} = 0$$

Gleichungen für die Materie und genauso wie in Übung 8.2 ergeben sich die Euler-Gleichungen

$$\begin{aligned} u^a \nabla_a \rho + (\rho + p) \nabla_a u^a &= 0, \\ (\rho + p) u^b \nabla_b u^a + (g^{ab} - u^a u^b) \nabla_b p &= 0, \end{aligned} \quad (9.4.3)$$

die hier allerdings etwas anders aussehen, wenn wir sie auf den vorliegenden Fall spezialisieren.

Betrachten wir zunächst die Gleichung

$$(\rho + p) u^b \nabla_b u^a + (g^{ab} - u^a u^b) \nabla_b p = 0.$$

Im zweiten Term $(g^{ab} - u^a u^b) \nabla_b p$ steht die inverse induzierte Metrik $g^{ab} - u^a u^b = 1/R^2 h^{ab}$. Diese enthält nur räumliche Richtungen, d.h., auf die Druckfunktion wirken nur räumliche Ableitungen. Da diese aber nur von t abhängt, verschwindet der gesamte zweite Term in der Gleichung und es bleibt nur noch die Geodätengleichung

$$u^b \nabla_b u^a = 0$$

übrig. Diese ist aber unter den gegebenen Umständen identisch erfüllt. Dies folgt aus der Übung 9.10, denn

$$u^b \nabla_b u_a = u^b \nabla_a u_b = \frac{1}{2} \nabla_a (u^b u_b) = 0.$$

Es verbleibt uns daher nur die Gleichung

$$u^a \nabla_a \rho + (\rho + p) \nabla_a u^a = 0.$$

Beachten wir, dass $u^a \nabla_a$ die Richtungsableitung in Richtung u^a bezeichnet, dann lässt sich die Gleichung in der Form

$$\dot{\rho} + (\rho + p) \nabla_a u^a = 0 \quad (9.4.4)$$

schreiben, wobei der Punkt die Ableitung nach t bedeutet. Wir müssen also die Divergenz $\nabla_a u^a$ der 4-Geschwindigkeit berechnen. Benutzen wir wieder die Übung 9.10, so sehen wir, dass

$$\nabla_a u^a = \nabla^a \nabla_a t = \square t$$

gilt. Wir brauchen also einen Ausdruck für den **d'Alembert-Operator** $\square$.

d'Alembert
Operator

Übung 9.11: Zeigen Sie, dass für eine beliebige Metrik

$$g = \sum_{i,k} g_{ik} dx^i dx^k$$

der d'Alembert-Operator durch den Ausdruck

$$\square f = \frac{1}{\sqrt{|g|}} \sum_{i,k} \partial_i \left(\sqrt{|g|} g^{ik} \partial_k f \right)$$

gegeben ist. Dabei bezeichnet $|g|$ den Absolutbetrag der Determinante der g_{ik} .

Verwenden wir diese Formel für $\square$ bei der Berechnung der Divergenz von u^a , so erhalten wir

$$\nabla_a u^a = \frac{\partial_t R^3}{R^3}$$

und aus (9.4.4) wird

$$\dot{\rho} + (\rho + p) \frac{\partial_t R^3}{R^3} = 0.$$

Diese Gleichung lässt sich auch in etwas anderer Form schreiben,

$$\partial_t (\rho R^3) = -p \partial_t (R^3) \quad (9.4.5)$$

oder, wenn wir R selbst als unabhängige Variable betrachten

$$\frac{\partial}{\partial R} (\rho R^3) = -3p R^2.$$

Die Interpretation dieser Gleichungen liegt nahe. Wir betrachten ein 3-dimensionales 'mitbewegtes Volumen' $V(t)$. Das ist ein Raumzeitbereich, in den isotrope Beobachter weder eintreten noch aus ihm austreten. Anschaulich gesprochen handelt es sich um eine feste Auswahl von Galaxien, die zeitlich verfolgt wird. Ähnlich wie sich der Abstand zweier isotroper Beobachter nur durch den Einfluss des Skalenfaktors verändert, so ändert sich auch der Rauminhalt des mitbewegten Gebiets nur aufgrund des Skalenfaktors. Das heißt, das Volumen $V(t)$ zum Zeitpunkt t ist

$$V(t) = R(t)^3 V_0,$$

wobei V_0 der mit der Metrik h gemessene, zeitlich unveränderliche Rauminhalt des Gebiets ist. Die Masse (Energie), die sich in diesem Volumen befindet, ist $M(t) = \rho(t)V(t)$, so dass aus der Gleichung (9.4.5) folgt

$$\frac{dM}{dt} = V_0 \frac{d(\rho R^3)}{dt} = -p V_0 \frac{dR^3}{dt} = -p \frac{dV}{dt}.$$

Diese Gleichung beschreibt die Änderung der Energie innerhalb des mitbewegten Volumens durch die an dem Volumen gegen den Druck p geleistete Arbeit.

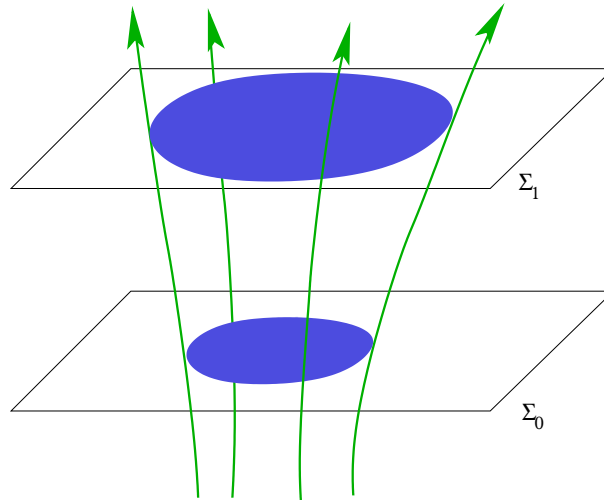


Abbildung 9.3: Ein mitbewegtes Volumen

Bevor wir mit den Einsteingleichungen die Dynamik des Skalenfaktors $R(t)$ bestimmen, wollen wir noch einige Betrachtungen zur Beobachtung in der Kosmologie anstellen. ‘Beobachten’ heisst, Signale, die aus dem Universum zu uns kommen, empfangen und auswerten. Signale laufen aber mit maximal Lichtgeschwindigkeit auf uns zu. Wenn wir also wissen wollen, welchen Weg diese Signale durch das Universum bis zu uns durchlaufen, dann müssen wir unseren Lichtkegel kennen, also die lichtartigen Geodäten bestimmen.

Dazu müssen wir die Geodätengleichung für Lichtstrahlen aufstellen. Wir stellen uns also vor, ein isotroper Beobachter zu sein, der um sich herum ein kugelsymmetrisches Universum wahrnimmt (Isotropie!). Dann kommen alle Signale aus radialer Richtung auf uns zu und wir können die Geodäten in der Form

$$\gamma(\lambda) = (t(\lambda), r(\lambda))$$

parametrisieren. Da es sich um lichtartige Geodäten handelt, gilt die Gleichung

$$0 = g(\dot{\gamma}, \dot{\gamma}) = \dot{t}^2 - \frac{R(t)^2 \dot{r}^2}{1 - kr^2}$$

oder

$$\frac{\dot{t}}{R(t)} = \pm \frac{\dot{r}}{\sqrt{1 - kr^2}}.$$

Für Geodäten, die bei $t = t_1$ im Abstand r ausgesandt werden und aus der Vergangenheit auf uns zu kommend zur Zeit $t = t_0$ bei uns ($r = 0$) eintreffen, ist $\dot{t} > 0$ und $\dot{r} < 0$, so dass wir die Gleichung

$$\int_{t_1}^{t_0} \frac{dt}{R(t)} = \int_0^r \frac{dr}{\sqrt{1 - kr^2}} \quad (9.4.6)$$

betrachten müssen. Diese Gleichung liefert uns die **kosmische Rotverschiebung**. Dass ein solcher Effekt auftreten muss, ist klar, wenn man bedenkt, dass die isotropen Beobachter relativ zueinander bewegt sind. Um ihn zu bestimmen, betrachten wir folgende Situation (vgl. Abb. 9.4). Ein isotroper Beobachter P (wir) empfängt Lichtsignale zum

kosmische
Rotverschi

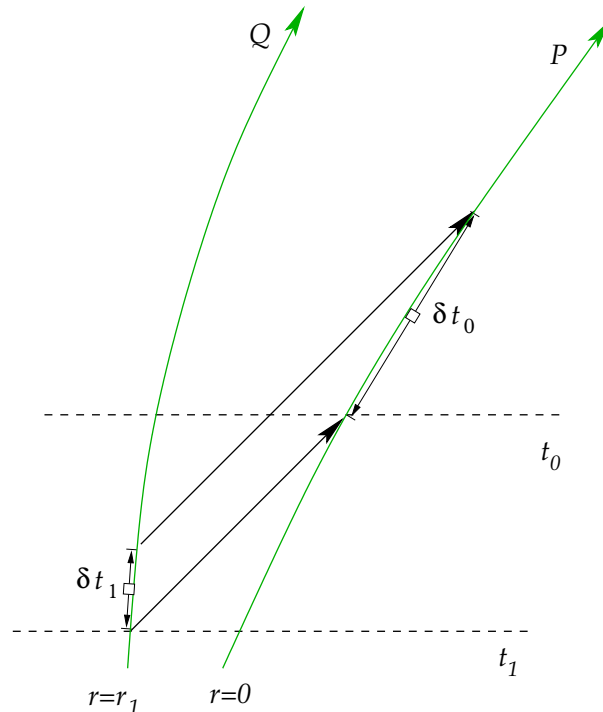


Abbildung 9.4: Zur kosmologischen Rotverschiebung

Zeitpunkt $t = t_0$ bzw. $t = t_0 + \delta t_0$, die ein zweiter Beobachter Q zum Zeitpunkt t_1 bzw. $t_1 + \delta t_1$ hintereinander ausgesendet hat. Für das frühere Signal gilt die Gleichung

$$\int_{t_1}^{t_0} \frac{dt}{R(t)} = \int_0^r \frac{dr}{\sqrt{1 - kr^2}}$$

und für das spätere Signal gilt entsprechend

$$\int_{t_1 + \delta t_1}^{t_0 + \delta t_0} \frac{dt}{R(t)} = \int_0^r \frac{dr}{\sqrt{1 - kr^2}}.$$

Weil der r -abhängige Teil für beide Signale der gleiche ist, folgt also

$$\int_{t_1}^{t_0} \frac{dt}{R(t)} = \int_{t_1 + \delta t_1}^{t_0 + \delta t_0} \frac{dt}{R(t)}.$$

Daraus ergibt sich

$$\int_{t_0}^{t_0 + \delta t_0} \frac{dt}{R(t)} = \int_{t_1}^{t_1 + \delta t_1} \frac{dt}{R(t)}.$$

Wenn wir δt als die Periode (und damit wegen $c = 1$ auch als die Wellenlänge) der gesendeten bzw. empfangenen Signale interpretieren, dann bleibt der Skalenfaktor innerhalb einer solchen Periode mit hoher Genauigkeit konstant, so dass wir schließlich die Gleichung

$$\frac{\delta t_0}{R(t_0)} = \frac{\delta t_1}{R(t_1)}$$

bekommen. Diese beschreibt das Phänomen der kosmischen Zeitdilatation bzw. die **kosmische Rotverschiebung**. Wenn Q ein Signal mit Wellenlänge $\lambda_1 = \delta t_1$ aussendet, dann empfängt P ein Signal mit der Wellenlänge

kosmische
Rotverschiebung

$$\lambda_1 = \lambda_0 \frac{R(t_1)}{R(t_0)}.$$

Man definiert die Rotverschiebung des Signals als die *relative Änderung der Wellenlänge* bezogen auf die gesendete Wellenlänge

$$z = \frac{\lambda_0 - \lambda_1}{\lambda_1} = \frac{R(t_0) - R(t_1)}{R(t_1)}.$$

In einem expandierenden Universum wie dem unseren ist $R(t_1) < R(t_0)$, so dass die Signale, die bei uns ankommen, einer Rotverschiebung $z > 0$ unterliegen.

Wenn die beiden Beobachter nahe beieinander sind, in dem Sinne, dass die Zeitdifferenz $\Delta t = t_0 - t_1$ klein gegenüber t_0 ist, dann ergibt sich

$$1 + z = \frac{R(t_0)}{R(t_1)} = \frac{R(t_0)}{R(t_0 - \Delta t)} \approx \frac{R(t_0)}{R(t_0) - \Delta t \dot{R}(t_0)} \approx 1 + \frac{\dot{R}(t_0)}{R(t_0)} \Delta t.$$

Andererseits ist aber auch

$$\int_{t_1}^{t_0} \frac{dt}{R(t)} = \int_{t_0 - \Delta}^{t_0} \frac{dt}{R(t)} \approx \frac{\Delta t}{R(t_0)}$$

und mit (9.4.6) folgt

$$\int_{t_1}^{t_0} \frac{dt}{R(t)} = \int_0^{r_1} \frac{dr}{\sqrt{1 - kr^2}} \approx r_1.$$

Fügen wir diese drei Gleichungen zusammen, ergibt sich die Relation

$$z \approx \dot{R}(t_0) r_1. \quad (9.4.7)$$

Die Rotverschiebung eines Signals, das von einer Galaxie ausgesendet wurde, ist also zu jeder Epoche proportional zur radialen Koordinate der Galaxie. Die Rotverschiebung kann als ein Maß für die Fluchtgeschwindigkeit der Galaxien betrachtet werden. Mit dieser Interpretation stellt diese Relation eine Proportionalität zwischen der Fluchtgeschwindigkeit der Galaxien und ihrem 'Abstand' von uns her.

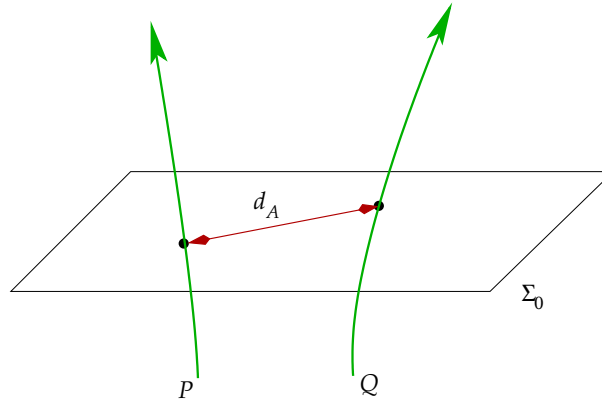


Abbildung 9.5: Zur Definition des absoluten Abstands

Das Problem, welches nun auftritt, ist die Unbeobachtbarkeit der Radialkoordinate r . Diese ist als Koordinate ohne unmittelbare physikalische Bedeutung. Wir müssen also die Frage untersuchen, wie man Distanzen im kosmologischen Kontext messen kann. Mathematisch betrachtet, ist es einfach, einen **absoluten Abstand** zwischen zwei Beobachtern (Galaxien) zu einer Zeit t_0 zu definieren (vgl. Abb. 9.5). Man bestimmt die Länge der Geodäte zwischen den Durchstosspunkten der Weltlinien von P (im Ursprung bei $r = 0$ angenommen) bzw. Q , erhält also

absoluten

$$d_A(P, Q, t_0) = R(t_0) \int_0^r \frac{dr}{\sqrt{1 - kr^2}}.$$

Dieser Abstand ist aber ebenfalls nicht beobachtbar, denn wir kennen den Wert von $R(t_0)$ nicht und haben keine Möglichkeit, über kosmische Dimensionen 'gleichzeitig' zu messen. Daher muss man andere Möglichkeiten in Betracht ziehen.

scheinbare Abstand

Eine solche Alternative ist der **scheinbare Abstand**. Dieser wird wie folgt definiert. Der Rückwärtskegel eines Beobachters P bei t_0 schneidet die Homogenitätshyperfläche Σ_1 in einer Kugeloberfläche von Punkten mit konstanter Koordinate $r = r_1$. Die Fläche dieser Kugel ist daher $4\pi R(t_1)^2 r_1^2$. Ein Objekt (einen Galaxienhaufen o.ä.) nimmt einen Anteil A an dieser Fläche ein, die für den Beobachter P einen Raumwinkel $\delta\omega$ abdeckt. Wären wir im Minkowski-Raum, dann würde zwischen der Fläche A und dem abgedeckten Raumwinkel die Beziehung $A = \delta\omega L^2$ bestehen, wobei L der Abstand zwischen Beobachter und Objekt ist. Wir *definieren* nun im gekrümmten Fall den scheinbaren Abstand durch eben diese Relation

$$d_S^2 = \frac{A}{\delta\omega}.$$

Nun ist aber

$$\frac{\delta\omega}{4\pi} = \frac{A}{4\pi R(t_1)^2 r_1^2},$$

so dass sich

$$d_S = R(t_1)r_1$$

ergibt. Wenn wir also ein Objekt beobachten, dessen tatsächliche Grösse wir kennen²,

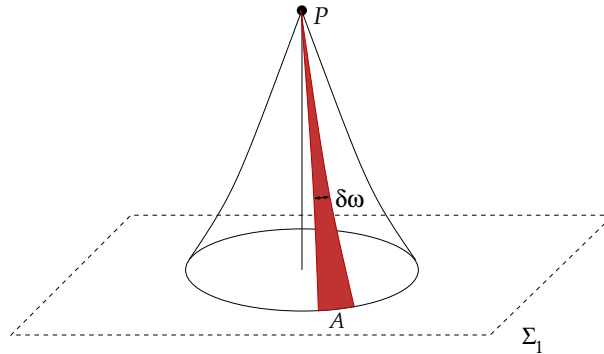


Abbildung 9.6: Zum scheinbaren Abstand

dann wissen wir auch die tatsächliche Fläche A und kennen nach Messung des Raumwinkels den scheinbaren Abstand des Objekts.

Eine weitere Alternative ist der *Helligkeitsabstand* d_H . Die Definition dieses Abstandsbegriffs beruht auf einer an sich einfachen Überlegung, die wir zuerst im flachen Newtonschen Raum durchführen. Eine Galaxie im Abstand L von uns sendet innerhalb eines Zeitintervalls δt_1 eine Anzahl N_1 von Photonen aus. Ein Teil davon, nämlich N_0 Stück werden von uns innerhalb des Zeitintervalls δt_0 auf einem Schirm der Fläche A aufgefangen. Nun ist im Newtonschen Fall $\delta t_0 = \delta t_1 = \delta t$ und wir bekommen die Beziehung

$$N_0 = \frac{A}{4\pi L^2} N_1.$$

Wenn wir also wissen, wieviele Photonen in der Galaxie innerhalb der Zeit δt_1 produziert werden und die ankommenden Photonen zählen, können wir via

$$L^2 = \frac{AN_1}{4\pi N_0}$$

den Abstand bestimmen. Um besser messbare Größen zu bekommen, erweitern wir den Bruch mit $h\nu/\delta t$ und erhalten

$$L^2 = \frac{A \frac{h\nu N_1}{\delta t}}{4\pi \frac{h\nu N_0}{\delta t}} = \frac{\frac{h\nu N_1}{\delta t}}{4\pi \frac{h\nu N_0}{A\delta t}}.$$

Der Zähler in diesem Ausdruck ist die abgestrahlte Strahlungsleistung P (Einheit Watt), während der Zähler (bis auf den Faktor 4π) die Intensität (Einheit Watt/qm) der Strahlung ist, die bei uns eintrifft. Kennen wir also die Strahlungsleistung der Quelle, dann liefert uns die Messung der ankommenden Strahlungsintensität den Abstand.

²Weil ähnliche Objekte sich auch viel näher bei uns befinden und dort mit direkteren Methoden ausgemessen werden können.

Im gekrümmten Fall sind einige Dinge anders: zunächst sind Emissions- und Detektionszeitintervall wegen der kosmischen Expansion nicht mehr gleich. Aus dem gleichen Grund sind auch die Frequenz der emittierten Photonen (ν_1) und die der empfangenen Photonen (ν_0) nicht mehr gleich. Vielmehr bestehen die Beziehungen

$$\frac{\delta t_0}{\delta t_1} = 1 + z = \frac{\nu_1}{\nu_0}$$

Die Intensität der empfangenen Strahlung ist daher

$$I = \frac{h\nu_0 N_0}{A \delta t_0}$$

und die Strahlungsleistung ist

$$P = \frac{h\nu_1 N_1}{\delta t_1}.$$

Helligkeitsabstand Definieren wir nun den **Helligkeitsabstand** wie im flachen Fall

$$d_H^2 = \frac{AN_1}{4\pi N_0}$$

und drücken die Photonenzahlen durch Leistung bzw. Intensität aus, dann bekommen wir

$$d_H^2 = \frac{1}{4\pi} \frac{P}{I} \frac{\delta t_1 / \nu_1}{\delta t_0 / \nu_0}$$

und damit die endgültige Definition des Helligkeitsabstands in der Form

$$d_H^2 = \frac{1}{4\pi} \frac{P}{I} \frac{1}{(1+z)^2}.$$

Die Bedeutung dieses Abstandsbegriffs liegt darin, dass man bei Kenntnis der Leuchtkraft von Objekten durch Messung von Intensität und Rotverschiebung ein Maß für die Distanz des Objekts bekommt. Die obige Definition ist im wesentlichen die im praktischen astronomischen Alltag verwendete, auch wenn diese noch etwas komplizierter ist. Um ein möglichst gutes und allgemein gültiges Maß für den Abstand zu bekommen, benötigt man Objekte, deren Leuchtkraft möglichst gut bekannt und einheitlich ist. Als solche 'Standardkerzen' werden Supernovae verwendet, weil man den Mechanismus der zur Supernova-Explosion führt, mittlerweile gut genug kennt und die Energieabstrahlung berechnen kann.

Mit diesem Abstandsbegriff sind wir in der Lage, eine Beziehung zwischen den beiden beobachtbaren Größen 'Rotverschiebung' und 'Helligkeitsabstand' herzustellen. Dazu bestimmen wir zuerst den Helligkeitsabstand einer Galaxie in Abhängigkeit von ihrer räumlichen Koordinate (vgl. Abb. 9.7) Wir, im Ereignis O_0 mit den Koordinaten $t = t_0$ und $r = 0$, empfangen Strahlung von der Galaxie in P_1 ($t = t_1$, $r = r_1$), die zum Zeitpunkt $t = t_1$ ausgesandt wurde. Diese Strahlung hat sich bis zur Zeit $t = t_0$ über die Oberfläche einer Kugel mit dem Zentrum in P_0 ($t = t_0$, $r = r_1$) ausgebreitet. Wegen

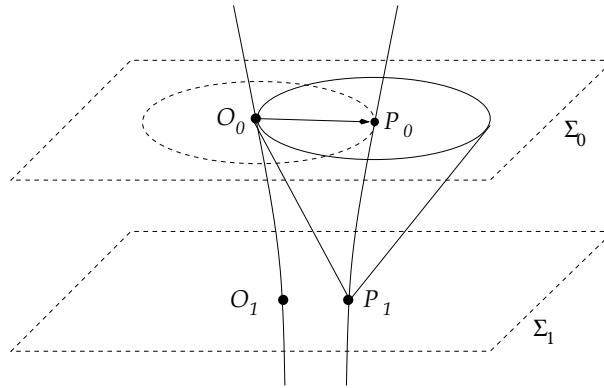


Abbildung 9.7: Zur Geometrie des Helligkeitsabstands

der Homogenität ist die Fläche dieser Kugel dieselbe, wie die Fläche der Kugel um O_0 , die durch P_0 geht. Die radiale Koordinate von P_0 ist $r = r_1$, so dass sich für den Flächeninhalt der Kugel

$$A = 4\pi(R(t_0)r_1)^2$$

ergibt. Wenn die Leistung der Quelle P ist, so ergibt sich beim Empfang der Strahlung eine Intensität von

$$I = \frac{P}{A(1+z)^2} = \frac{P}{4\pi(R(t_0)r_1)^2(1+z)^2},$$

wenn man die Rotverschiebung mit berücksichtigt. Aus dieser Relation und der Definition des Helligkeitsabstands ergibt sich die Abhängigkeit des Helligkeitsabstands von der Radialkoordinate und dem Skalenfaktor in der Form

$$d_H = R(t_0)r_1.$$

Damit können wir nun Rotverschiebung und Helligkeitsabstand zueinander in Beziehung setzen. Mit (9.4.7) erhalten wir nach Elimination der unbeobachtbaren Radialkoordinate

$$z \approx \frac{\dot{R}(t_0)}{R(t_0)} d_H = H(t_0) d_H, \quad (9.4.8)$$

wobei wir hier den **Hubble-Parameter** $H(t_0) = \dot{R}(t_0)/R(t_0)$ eingeführt haben³. Wir bekommen also aus diesen Überlegungen das bekannte **Hubblesche Gesetz**, wonach die Rotverschiebung von 'nahen' Galaxien(haufen), und damit ihre Fluchtgeschwindigkeit, proportional zu ihrem Abstand zu uns ist. Die Einheit des Hubble-Parameters ist $1/s$, also inverse Zeit. Er wird aber immer in der Einheit $(\text{km/s})/\text{Mpc}$ angegeben. Der Wert der Hubble-Konstanten (also der Wert des Hubbleparameters zur heutigen Zeit) war lange Zeit heftig umstritten. Die Werte lagen je nach Arbeitsgruppe zwischen 50 und 100 $(\text{km/s})/\text{Mpc}$. Nach Messungen des Satelliten WMAP scheint der Wert

Hubble-Parameter
Hubblesche Gesetz

$$71 \pm 4 (\text{km/s})/\text{Mpc} \approx 24 \times 10^{-19} 1/s.$$

³Offensichtlich ist dies keine Konstante, sondern eine zeitabhängige Größe.

heute jedoch realistisch zu sein.

Schließlich betrachten wir noch eine weitere Folgerung aus der Gleichung für den Rückwärtslichtkegel eines Beobachters O , der sich zur Zeit $t = t_0$ im Ereignis O_0 befindet. Der Rückwärtslichtkegel besteht aus allen Ereignissen P mit Koordinaten (t, r) ist, für die die Gleichung

$$\int_t^{t_0} \frac{dt'}{R(t')} = \int_0^r \frac{dr'}{\sqrt{1 - k(r')^2}} \quad (9.4.9)$$

gilt. Für den Vorwärtslichtkegel des Ereignisses O_1 (vgl. Abb. 9.8) gilt dann die ganz

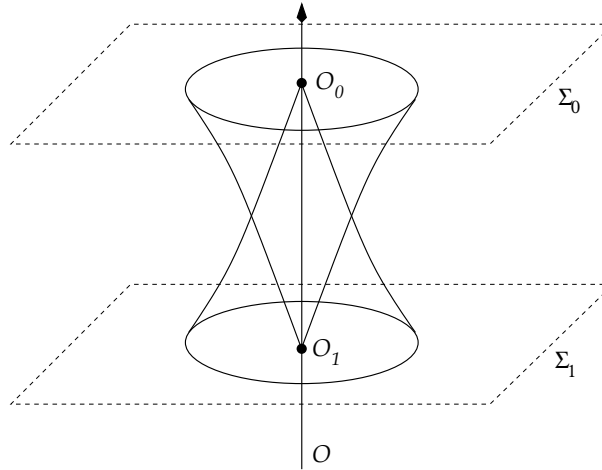


Abbildung 9.8: Vorwärts- und Rückwärtslichtkegel von Ereignissen auf der Weltlinie eines Beobachters O

entsprechende Gleichung

$$\int_{t_1}^t \frac{dt'}{R(t')} = \int_0^r \frac{dr'}{\sqrt{1 - k(r')^2}} \quad (9.4.10)$$

Abhängig von der konkreten Form des Skalenfaktors $R(t)$ kann es vorkommen, dass die linke Seite der Gleichung einen endlichen Wert behält, selbst wenn t_0 beliebig groß wird, wenn also

$$\lim_{t_0 \rightarrow \infty} \int_t^{t_0} \frac{dt'}{R(t')} = L < \infty$$

ist. Dies bedeutet, dass es einen maximalen absoluten Abstand gibt aus dem Licht zum Beobachter innerhalb seiner Existenz gelangen kann. Das bedeutet, es gibt Ereignisse, die der Beobachter nie sehen wird, egal wie lange er wartet. Man sagt, es gibt einen *kosmologischen Ereignis-Horizont* für den Beobachter O . Dieser ist demnach eine Trennung zwischen Ereignissen, die der Beobachter noch sehen kann und solchen, die ihm für immer verborgen sind, also hinter dem Horizont liegen.

Wir erinnern uns an den Ereignishorizont, der in der Schwarzschild-Raumzeit auftrat. Hier wie dort ist der Horizont eine *lichtartige* Hyperfläche, die nur in einer Richtung

Ereignis-Horizont

durchlaufen werden kann. Anders als der Schwarzschild-Horizont ist der kosmologische Ereignis-Horizont jedoch an Beobachter gekoppelt. Die Ereignis-Horizonte verschiedener Beobachter sind verschieden während der Ereignis-Horizont der Schwarzschild-Raumzeit für alle äußeren Beobachter derselbe ist.

Eine weitere Sorte von Horizonten ergibt sich, wenn man den zeitumgekehrten Fall betrachtet. Der Vorwärtslichtkegel des Ereignisses O_1 besteht aus allen Ereignissen, die dieses Ereignis 'sehen'. Sollte es nun der Fall sein, dass die linke Seite in (9.4.10) einen endlichen Wert annimmt, wenn t_1 in die Vergangenheit verschoben wird (entweder in die unendliche Vergangenheit $t_1 \rightarrow \infty$ oder bis zu einem endlichen Anfangswert, z.B. $t_1 = 0$)

$$\lim_{t_1 \rightarrow -\infty} \int_{t_1}^t \frac{dt'}{R(t')} = L < \infty$$

gilt, dann gibt es Ereignisse im Universum, die von der Existenz des Beobachters O nie Kenntnis erlangen werden. Auch hier bildet die Trennfläche zwischen Ereignissen, die O kennen, und solchen, die O nicht kennen, eine lichtartige Hyperfläche. Man nennt sie den **Teilchenhorizont** von O . Der Teilchenhorizont besteht aus all jenen Ereignissen, die von Photonen durchlaufen werden, die zum Beginn des Universums vom Beobachter O ausgesendet wurden.

Teilchenhorizont

9.5 Die Dynamik des Universums

Nach diesen kinematischen Betrachtungen, die ohne die explizite Lösung der Einstein-Gleichung auskommen, wenden wir uns nun den Feldgleichungen zu, um die möglichen Skalenfaktoren $R(t)$ zu bestimmen. Dazu müssen wir die Komponenten der Gleichung

$$G_{ab} - \Lambda g_{ab} = -8\pi T_{ab}$$

bestimmen⁴ Dies sind zunächst 10 Gleichungen, denn es handelt sich um eine Gleichung zwischen symmetrischen Tensoren 2. Stufe in 4 Dimensionen. Nun kann man aber leicht sehen, dass viele dieser Gleichungen aufgrund der vorausgesetzten Homogenität und Isotropie trivial werden müssen.

Übung 9.12: Schreiben Sie $E_{ab} = G_{ab} - \Lambda g_{ab} + 8\pi T_{ab}$, so dass die Feldgleichungen einfach durch $E_{ab} = 0$ gegeben sind. Zeigen Sie, dass für die Friedmann-Robertson-Walker-Metrik (9.4.1) folgendes gilt

1. Die Zeit-Raum-Komponente von E_{ab} muss identisch verschwinden

$$E_{ti} = 0 \quad \text{für } i = 1, 2, 3.$$

2. Die Raum-Raum-Komponenten E^i_k mit $i, k = 1, 2, 3$ sind reine Spur, d.h. $E^i_k = E \delta^i_k$ für eine *räumliche Konstante* E .

⁴Aus aktuellem Anlass berücksichtigen wir hier auch die kosmologische Konstante Λ .

Folgern Sie daraus, dass nur zwei Gleichungen von den 10 Gleichungen $E_{ik} = 0$ nicht identisch erfüllt sind.

Die einzigen nichttrivialen unabhängigen Komponenten der Feldgleichungen sind die 00- und die 11-Komponente, aus denen wir nach etwas Umformen die beiden Gleichungen

$$3\frac{\dot{R}^2}{R^2} + 3\frac{k}{R^2} - \Lambda = 8\pi\rho, \quad (9.5.1)$$

$$2\frac{\ddot{R}}{R} + \frac{\dot{R}^2}{R^2} + \frac{k}{R^2} - \Lambda = -8\pi p. \quad (9.5.2)$$

Die erste dieser Gleichungen lässt sich auch in der Form

$$\dot{R}^2 + k - \frac{1}{3}\Lambda R^2 = \frac{8}{3}\pi\rho R^2 \quad (9.5.3)$$

Friedmann-Gleichung schreiben. Diese Gleichung nennt man die **Friedmann-Gleichung**. Sie entspricht der Erhaltungsgleichung für die Energie in der Newtonschen Mechanik. Dies erkennt man am einfachsten, indem man sie in der Form

$$\frac{1}{2}\dot{R}^2 - \frac{4\pi}{3}\left(\rho + \frac{\Lambda}{8\pi}\right)R^2 = -\frac{1}{2}k = \text{const.}$$

schreibt. Der erste Term entspricht der kinetischen Energie, während der zweite Term die potentielle Energie repräsentiert. Hierbei spielt nicht nur die gravitative Bindungsenergie der Materie eine Rolle, sondern auch die von der kosmologischen Konstanten verursachte Energie trägt zur Gesamtbilanz bei.

Übung 9.13: Zeigen Sie, dass aus (9.5.1) und (9.5.2) die Erhaltungsgleichung (9.4.5) folgt.

Wir können also anstelle der Feldgleichungen (9.5.1) und (9.5.2) die Friedmann-Gleichung und die Erhaltungsgleichung als bestimmende Gleichungen benutzen. Wir haben also die Gleichungen

$$\begin{aligned} \dot{R}^2 + k - \frac{1}{3}\Lambda R^2 &= \frac{8}{3}\pi\rho R^2 \\ \partial_t(\rho R^3) &= -p\partial_t(R^3) \end{aligned}$$

zu lösen. Die Unbekannten dabei sind der Skalenfaktor $R(t)$, sowie die Druckfunktion $p(t)$ und die Dichtefunktion $\rho(t)$. Wir haben also zwei Gleichungen für drei Funktionen. Die fehlende Information ist genauso wie im Kapitel 8 in der *Zustandsgleichung* zu suchen, die Druck und Dichte miteinander verknüpft und damit festlegt, um welche Materie es sich in dem jeweils betrachteten Modell handelt. Es handelt sich dabei um eine *unabhängige* Annahme über das Modell, die nichts mit den Annahmen über die Geometrie zu tun hat.

Die verschiedenen kosmologischen Modelle werden nun also bestimmt durch die Wahl der Konstanten k , die festlegt ob wir eine sphärische, flache oder hyperbolische räumliche Geometrie vor uns haben, vom Vorzeichen der kosmologischen Konstanten und

vom konkreten Materiemodell. Dies ergibt eine große Auswahl von Möglichkeiten, die wir hier nicht alle diskutieren können. Daher beschränken wir uns auf einige spezielle und einfach zu berechnende Fälle.

9.5.1 Kosmologische Modelle mit $\Lambda = 0$

Staubmaterie

Wir betrachten zuerst den Fall verschwindender kosmologischer Konstante $\Lambda = 0$. Nehmen wir zunächst an, das Universum sei mit Staub, das heißt mit *inkohärenter Materie* gefüllt. Damit sind Teilchen gemeint, die massebehaftet sind, und nur gravitativ miteinander in Wechselwirkung stehen. Das bedeutet, dass der Druck in dieser Materieverteilung verschwinden muss. Die Zustandsgleichung ist daher

$$p = 0.$$

Nun folgt aus der Erhaltungsgleichung sofort

$$\partial_t (\rho R^3) = 0$$

und damit folgt die Existenz einer Konstanten M , so dass

$$\rho(t) = \frac{M}{\frac{4\pi}{3}R(t)^3}.$$

Setzen wir dies in die Friedmann-Gleichung (mit $\Lambda = 0$) ein, bekommen wir

$$\dot{R}^2 = \frac{2M}{R} - k.$$

Offensichtlich hängt die Lösung dieser Gleichung von k ab und wir diskutieren daher die drei Fälle nacheinander.

Der Fall $k = 0$ Hier ergibt sich die Gleichung (wir wählen das positive Vorzeichen weil wir $\dot{R} > 0$ annehmen)

$$R\dot{R}^2 = 2M \iff \sqrt{R}\dot{R} = +\sqrt{2M}$$

mit der Lösung

$$R(t) = R_0 \left(\frac{3}{2} + \frac{3}{2} \sqrt{\frac{2M}{R_0}} \frac{(t - t_0)}{R_0} \right)^{\frac{2}{3}}.$$

Dabei ist $R_0 = R(t_0)$. Wir sehen, dass der Skalenfaktor für ein bestimmtes t verschwindet, und wählen die freie Konstante t_0 oBdA als eben dieses t . Dann gilt $R(0) = 0$ und wir erhalten die Lösung

$$R(t) = \left(\frac{3}{2} \sqrt{2M} t \right)^{\frac{2}{3}}.$$

Damit ergibt sich für die Dichte die Zeitabhängigkeit

$$\rho(t) = \frac{1}{6\pi t^2}.$$

In den Fällen $k = \pm 1$ haben wir die Gleichung

$$R\dot{R}^2 + kR = 2M$$

vorliegen. Hier bietet es sich an, eine neue Variable s über die Relation

$$\dot{s} = 1/R$$

einzuführen. Dann gilt

$$\left(\frac{dR}{ds} \right)^2 = R^2 \dot{R}^2 = 2M R - kR^2 \quad (9.5.4)$$

und damit auch

$$\frac{d^2 R}{ds^2} = -kR + M.$$

Dies ist eine lineare, inhomogene Gleichung für $R(s)$, deren allgemeine Lösung sich aus der Lösung der homogenen Gleichung und einer partikulären Lösung zusammensetzt. Wir schreiben daher

$$R(s) = Ae^{\omega s} + Be^{-\omega s} + kM.$$

Dabei ist $\omega^2 = -k$ gesetzt. Dies ist natürlich die Lösung der Differentialgleichung 2.Ordnung, die zwei beliebige Konstanten A und B enthält, während wir die Lösung der Gleichung 1. Ordnung (9.5.4) suchen, die nur eine willkürliche Konstante enthalten darf. Wir setzen daher die allgemeine Lösung in (9.5.4) ein und erhalten

$$(\omega Ae^{\omega s} - \omega Be^{-\omega s})^2 = 2M (Ae^{\omega s} + Be^{-\omega s} + kM) - k (Ae^{\omega s} + Be^{-\omega s} + kM)^2$$

Nach Ausmultiplizieren ist leicht zu sehen, dass die einzigen Terme, die übrig bleiben, auf die Beziehung

$$2kAB = -2kAB + M^2 \iff AB = \frac{M^2}{4}$$

führen. Wählen wir $A = B = kM/2$ (dies entspricht der Festlegung der willkürlichen Konstanten so, dass $R(0) = 0$ gilt), dann finden wir die Lösung

$$R(s) = \frac{kM}{2} (2 - e^{\omega s} - e^{-\omega s}).$$

Der Fall $k = 1$ Hier ist $k = 1$ und $\omega = i$ und die Lösung lässt sich schreiben als

$$R(s) = M (1 - \cos s).$$

Mit dieser Lösung ergibt sich für t die Gleichung

$$\frac{dt}{ds} = M (1 - \cos s)$$

mit der Lösung (mit $t(0) = 0$)

$$t(s) = M (s - \sin s).$$

Damit haben wir die Lösung $R(t)$ in parametrischer Form, d.h. als Paar $(R(s), t(s))$ gefunden. Die so beschriebene Kurve ist eine *Zykloide*.

Der Fall $k = -1$ Hier ist nun $k = -1$ und $\omega = 1$. Die Lösung lässt sich daher in der Form

$$R(s) = M (\cosh s - 1)$$

schreiben. Dann ergibt sich wiederum für t die Lösung

$$t(s) = M (\sinh s - s).$$

Auch hier haben wir eine parametrische Darstellung der Lösung $R(t)$. Das Verhalten der Lösung $R(t)$ in den drei Fällen ist in Abb. 9.9 verdeutlicht. Das Verhalten dieser drei

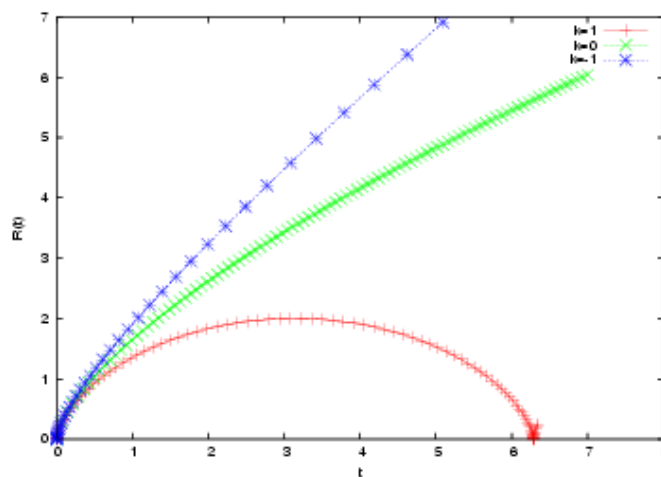


Abbildung 9.9: Der Skalenfaktor in den drei Fällen für Staubmaterie

Lösungen lässt sich einfach feststellen. Alle drei Lösungen starten für $t = 0$ mit $R = 0$ bei $t = 0$. Daher ist in allen drei Fällen die Dichte wegen $\rho \sim R^{-3}$ *unendlich groß*. Es handelt sich hier um eine physikalische Singularität, eine *Krümmungssingularität*. Im Fall

$k = 1$ hat der Skalenfaktor ein Maximum und nimmt dann wieder auf $R = 0$ ab. Das Universum dehnt sich bis zu einem maximalen Wert aus und kollabiert danach wieder auf eine Singularität zusammen. Im Gegensatz zu diesem *geschlossenen Universum* liefern die beiden anderen Fälle ein *offenes Universum*. Hier wächst der Skalenfaktor unbegrenzt an. Der Unterschied besteht in der Expansionsrate.

Übung 9.14: Zeigen Sie, dass im Fall $k = 0$ die Proportionalität $R \sim t^{2/3}$ besteht, während im hyperbolischen Fall ($k = -1$) $R \sim t$ gilt.

Es gibt für diese Klasse von Modellen neben dem diskreten Parameter k nur noch einen Parameter M , der kontinuierliche Werte annehmen kann.

Ein Strahlungsuniversum

Als weitere Modellklasse betrachten wir den Fall eines Universums, das mit Strahlung (Photonen) angefüllt ist. Für eine solche Strahlungsmaterie ist der Energie-Impuls-Tensor spurfrei, d.h., es ist $T_a^a = 0$.

Übung 9.15: Zeigen Sie, dass der Energie-Impuls-Tensor des elektromagnetischen Feldes spurfrei ist (vgl. Vorlesung zur SRT).

Für den Energie-Impuls-Tensor (9.4.2) eines homogenen, isotropen Universums folgt daraus sofort $\rho = 3p$. Dies ist die Zustandsgleichung, die wir für diesen Fall verwenden wollen.

Setzen wir diese Zustandsgleichung in die Kontinuitätsgleichung (9.4.5) ein, bekommen wir

$$\partial_t (\rho R^3) = -\rho R^2 \dot{R},$$

was sich in die Gleichung

$$\partial_t (\rho R^4) = 0$$

umformen lässt. Damit ergibt sich wiederum die Existenz einer räumlichen Konstanten M , so dass wir

$$\frac{4\pi}{3} \rho(t) = \frac{M}{R(t)^4}$$

erhalten. Damit liefert uns die Friedmann-Gleichung

$$\dot{R}^2 = \frac{2M}{R^2} - k.$$

Mit demselben Trick wie oben können wir auch diese Gleichung lösen. Wir führen wieder die Variable s wie oben ein und erhalten für $R(s)$ die Differentialgleichung

$$\frac{d^2 R}{ds^2} = -kR^2 + 2M.$$

Diese lässt sich nun für die drei Fälle $k = 0, \pm 1$ leicht lösen. Wir erhalten die Lösungen mit $R = 0$ für $t = 0$ in der folgenden Liste

(i) $k = 0$ (offen):

$$R(t) = (8M)^{1/4} \sqrt{t},$$

(ii) $k = 1$ (geschlossen):

$$R(s) = \sqrt{2M} \sin s, \quad t(s) = \sqrt{2M} (1 - \cos s),$$

(iii) $k = -1$ (offen)

$$R(s) = \sqrt{2M} \sinh s, \quad t(s) = \sqrt{2M} (\cosh s - 1).$$

Das Verhalten dieser Lösungen ist in Abb. 9.10 graphisch dargestellt. Das Verhalten

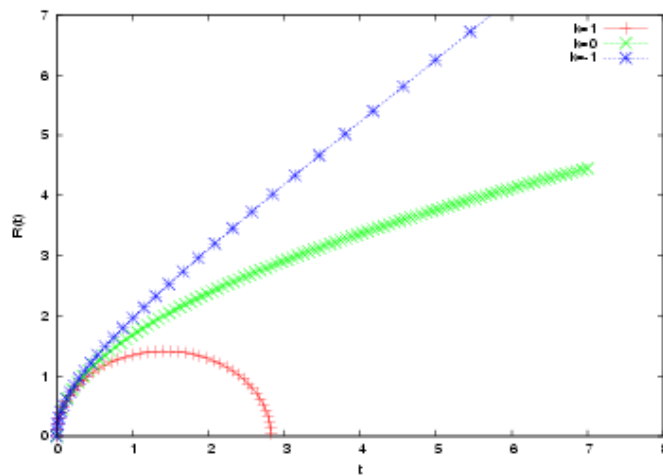


Abbildung 9.10: Der Skalenfaktor in den drei Fällen für Strahlungsmaterie

des Skalenfaktors ist qualitativ genau das gleiche wie für Staubmaterie. Jedoch ist das Verhalten der Expansionsrate am Anfang und am Ende verschieden. Auch ist die Lebensdauer des geschlossenen Universums anders.

Übung 9.16: Diskutieren Sie diese Unterschiede zwischen den Modellen.

Offensichtlich haben alle Modelle, die wir bislang untersucht haben eine Anfangssingularität. Dies ist notwendigerweise so. Denn wenn wir aus den Gleichungen (9.5.1) und (9.5.2) den Term mit $\dot{R}$ eliminieren, dann bekommen wir die Gleichung

$$2 \frac{\ddot{R}}{R} = -\frac{8\pi}{3}(\rho + 3p) + \frac{2}{3}\Lambda,$$

also im vorliegenden Fall mit $\Lambda = 0$

$$\frac{\ddot{R}}{R} = -\frac{4\pi}{3}(\rho + 3p).$$

Unabhängig von der Zustandsgleichung erwartet man, dass für jede Materie die Ungleichung

$$\rho + 3p > 0$$

gilt.

Übung 9.17: Zeigen Sie, dass diese Ungleichung aus der sogenannten *dominanten Energiebedingung* folgt: für jeden zeitartigen Vektor u^a ist $\rho^a = T^{ab}u_b$ ebenfalls zeitartig. Wie lässt sich diese Bedingung physikalisch interpretieren, wenn man weiß, dass ρ^a der Vektor der Energieflußdichte ist?

Unter dieser Annahme folgt, dass $\ddot{R} < 0$. Der Graph der Funktion $R(t)$ ist also konvex. Die Expansionsgeschwindigkeit nimmt also in positiver Zeitrichtung ab, bzw. die Kontraktionsgeschwindigkeit nimmt in negativer Zeitrichtung zu. Wenn es also einen Zeitpunkt gibt, an dem $\dot{R}$ positiv ist, dann gab es vorher einen Zeitpunkt, an dem $R = 0$ war. Ist t_0 dieser Zeitpunkt und ist $\dot{R}_0 = \dot{R}(t_0) > 0$ und $R_0 = R(t_0) > 0$, dann war der Zeitpunkt t_* an dem $R = 0$ verschwand frühestens vor einem Zeitintervall $t_0 - t_* = R_0/\dot{R}_0$. Das bedeutet, die Hubble-Konstante H_0 liefert uns eine Abschätzung für das Alter T des Universums. Mit den Zahlen aus dem vorigen Abschnitt ergibt sich

$$T \approx H_0^{-1} \approx 10^{10} \text{a.}$$

Anfangssingularität

Zusammengefasst ergibt sich vorerst das folgende Bild. Alle bisher betrachteten Lösungen beginnen mit einer **Anfangssingularität** eine Expansionsphase. Je nach Modell hält diese Phase unendlich lange an, nämlich wenn die Raumschnitte unendlich ausgedehnt sind, oder aber das Universum erreicht eine maximale Ausdehnung und kontrahiert dann wieder bis zu einer Singularität, wenn die Raumschnitte kompakt sind. Die exakten Werte der Expansionsrate, des Alters des Universums usw. hängen dabei vom Modell ab

9.5.2 Kosmologische Modelle mit $\Lambda \neq 0$

Wir kommen nun zu Modellen mit nichtverschwindender kosmologischer Konstanten. Wir fragen zuerst nach der Möglichkeit einer zeitunabhängigen Lösung der Gleichungen (9.5.1) und (9.5.2). Setzen also $\dot{R} = 0$ in diesen Gleichungen, bekommen wir

$$3k - \Lambda R^2 = 8\pi\rho R^2,$$

$$k - \Lambda R^2 = -8\pi p R^2.$$

Daraus folgt durch eliminieren der kosmologischen Konstante

$$k = 4\pi(\rho + p)R^2.$$

Beschränken wir uns auf realistische Materiemodelle, dann ist $\rho + p \geq 0$, d.h. $k = 0$ oder $k = 1$. Als Beispiel sei wieder ein Staubmodell zugrunde gelegt. Dann haben wir

$p = 0$ und es kann nur $k = 1$ zu einer nicht-trivialen Raumzeit führen (was würde $k = 0$ liefern?). Dann gilt also

$$4\pi\rho R^2 = 1, \quad \Lambda R^2 = 1, \quad k = 1$$

und die Metrik lautet

$$g = dt^2 - R^2 \left(d\chi^2 + \sin^2 \chi (d\theta^2 + \sin^2 \theta d\phi^2) \right). \quad (9.5.5)$$

Diese Metrik beschreibt das **Einsteinsche statische Universum**. Der einzige freie Parameter in dieser Lösung ist der Radius R des Universums, den wir auch noch als $R = 1$ setzen können, denn er bestimmt im wesentlichen die Größenskala, d.h. die Einheit in der wir Längen messen. Die Parameter in diesem Modell sind also exakt fixiert, um das Gleichgewicht zu halten und Zeitunabhängigkeit zu erreichen. Es ist ziemlich klar, dass eine Abweichung von diesen Werten ein Universum zur Folge hätte, welches sich von diesen Werten wegbewegen würde. Eine genauere Störungsrechnung zeigt, dass diese Lösung instabil ist.

Für zeitabhängige Modelle beschränken wir uns im Weiteren auf den Fall $k = 0$ und betrachten auch hier nur Staubmaterie, $p = 0$. Dann lautet die Friedmann-Gleichung

$$\dot{R}^2 = \frac{2M}{R} + \frac{1}{3}\Lambda R^2$$

mit der Konstanten $2M = \frac{8\pi}{3}\rho R^3$, deren Existenz aus der Energieerhaltung folgt.

Wir nehmen $\Lambda > 0$ an und setzen

$$u = \frac{\Lambda}{3M} R^3.$$

Dann folgt die Gleichung

$$\dot{u}^2 = 3\Lambda(u^2 + 2u),$$

die sich via

$$\frac{\dot{u}}{\sqrt{(u+1)^2 - 1}} = \sqrt{3\Lambda}$$

und der Substitution $u + 1 = \cosh x$ zu

$$u = \cosh(x_0 + \sqrt{3\Lambda}(t - t_0)) - 1$$

integrieren läßt. Setzen wir als Anfangsbedingung noch $R(0) = 0$ voraus, dann bekommen wir die Lösung

$$\frac{\Lambda}{3M} R^3 = \cosh(\sqrt{3\Lambda}t) - 1.$$

Im Falle $\Lambda < 0$ setzen wir

$$u = -\frac{\Lambda}{3M} R^3$$

und erhalten nach ganz analogem Vorgehen die Lösung

$$\frac{\Lambda}{3M} R^3 = 1 - \cos(\sqrt{-3\Lambda}t).$$

Wir erhalten aus diesen Lösungen auch den Grenzfall $\Lambda \rightarrow 0$, indem wir z.B. die Taylor-Entwicklung der Lösung mit positivem Λ aufschreiben

$$R^3 = \frac{3M}{\Lambda} \left(\cosh(\sqrt{3\Lambda}t) - 1 \right) \approx \frac{3M}{\Lambda} \left(1 + \frac{3}{2}\Lambda t^2 + \frac{3}{8}\Lambda^2 t^4 + \mathcal{O}(t^6) - 1 \right)$$

und dann den Grenzübergang $\Lambda \rightarrow 0$ durchführen. Es ergibt sich (natürlich) die oben erhaltene Lösung

$$R = \left(\frac{9}{2} M t^2 \right)^{\frac{1}{3}}.$$

Das Verhalten dieser Lösungen in Abb. 9.11 angegeben. Wir sehen auch hier wieder,

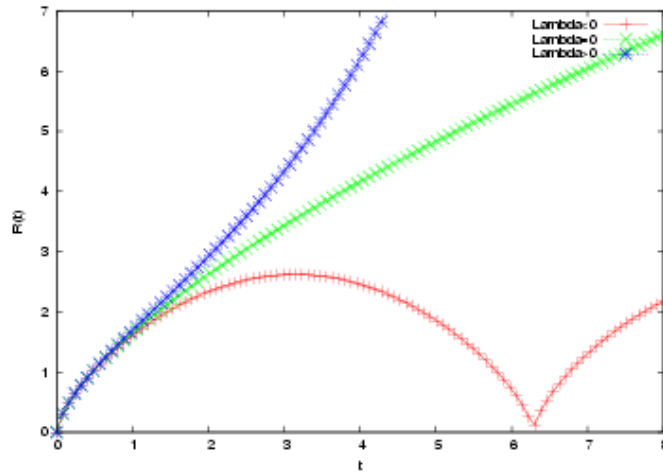


Abbildung 9.11: Die drei Möglichkeiten für Modelle flache Raumschnitte mit verschiedener kosmologischer Konstanten $\Lambda = \pm 1/3, 0$

dass es offene und geschlossene Universen gibt. Für *positives* Λ haben wir ein offenes Modell, während *negatives* Λ dazu beiträgt, das Universum zu schließen. Die Expansionsrate zu Anfang ist unabhängig von Λ und in allen Fällen gegeben durch

$$R \sim t^{\frac{2}{3}}.$$

Der Wert der kosmologischen Konstanten Λ beeinflusst jedoch die Expansionsrate des offenen Modells für große t , die zu

$$R \sim e^{\sqrt{\frac{\Lambda}{3}}t}$$

bestimmt wird, während die Lebenszeit des geschlossenen Modell ebenfalls durch Λ beeinflusst wird.

$$T = \frac{2\pi}{(-3\Lambda)}.$$

9.5.3 Qualitative Diskussion der Friedmann-Lösungen

Nach dieser (unvollständigen) Liste mit expliziten Ausdrücken für die Lösung der Friedmann-Gleichung in verschiedenen Fällen betrachten wir diese Gleichung zum Schluß noch einmal unter einem qualitativen Gesichtspunkt. Die Gleichung war

$$\dot{R}^2 + k - \frac{1}{3}\Lambda R^2 = \frac{8}{3}\pi\rho R^2,$$

und nimmt für Staubmaterie die Form

$$\dot{R}^2 = \frac{2M}{R} + \frac{1}{3}\Lambda R^2 - k =: F(R)$$

an. Die Form der Funktion $F(R)$ gibt uns schon viel Information über das qualitative Verhalten der Lösungen. Offensichtlich divergiert $F(R)$ für $R \rightarrow 0$ zu beliebig großen positiven Werten. Dies ist Ausdruck der Tatsache, dass das Universum mit einer Anfangssingularität beginnt.

Beginnen wir mit dem Fall negativer kosmologischer Konstante. In diesem Fall wird $F(R) \approx \frac{1}{3}\Lambda R^2 < 0$ für große Werte von R . Es gibt also eine Nullstelle von $F(R)$, eine Stelle an der $\dot{R}$ verschwinden muss. Folglich muss es sich bei allen diesen Modellen um ein geschlossenes Universum handeln.

Für $\Lambda = 0$ kennen wir die exakten Lösungen. Wir können ihr Verhalten aber auch aus der Betrachtung von $F(R)$ erschließen. Es gibt in diesem Fall nur für $k = 1$ eine Nullstelle von $F(R)$, die bei endlichen Radien liegt. Nur dieser Fall ist daher geschlossen. Die anderen Fälle sind offen. $k = 0$ liefert den Grenzfall, in dem die Expansionsrate gegen Null geht, wenn R beliebig groß wird. Der Fall $k = -1$ expandiert asymptotisch mit konstanter Geschwindigkeit.

Der komplizierteste Fall ist der mit positiver kosmologischer Konstante. Nun ist $F(R) \approx \frac{1}{3}\Lambda R^2 > 0$ für große Werte von R . Das heißt F hat ein Minimum und die Frage ist, wie groß ist der Wert von F an diesem Minimum? Denn wenn F dort negativ ist, dann gibt es ein geschlossenes Modell, ansonsten ein offenes. Die Ableitung von F ist

$$F'(R) = -\frac{2M}{R^2} + \frac{2}{3}\Lambda R$$

und dieser Ausdruck verschwindet für

$$R_* = \sqrt[3]{\frac{3M}{\Lambda}}.$$

Der Wert von F an diesem Punkt ist

$$F(R_*) = \frac{3M}{R_*} - k.$$

In den Fällen $k = 0$ und $k = -1$ ist dieser Wert offensichtlich positiv, so dass sich hier ein offenes Universum ergibt. Im Fall $k = 1$ jedoch gibt es Werte von Λ , für die $F(R_*)$ negativ werden kann. Wenn Λ den kritischen Wert

$$\Lambda_c = \frac{1}{9M^2}$$

annimmt, dann ist $F(R_*) = 0$. In diesem kritischen Fall sind drei verschiedene Modelle denkbar, die sich durch die Anfangsbedingungen unterscheiden:

- (a) Einsteins statisches Universum mit $R = R_*$
- (b) ein offenes Universum, das von einer Anfangssingularität ausgehend expandiert und sich asymptotisch dem statischen Universum nähert
- (c) ein offenes Universum, das für immer expandiert, aber in der Vergangenheit asymptotisch dem statischen Universum gleich war.

Im überkritischen Fall $\Lambda > \Lambda_c > 0$ gibt es keine Nullstelle. Es gibt daher nur ein offenes Modell, das Lemaître-Modell, welches aber die Besonderheit hat, dass auf eine Phase starker Expansion eine „ruhige“ Phase folgt, in der $\dot{R}$ klein ist und woran sich wieder eine Phase starker Expansion anschließt.

Im unterkritischen Fall $\Lambda_c > \Lambda > 0$ gibt es zwei Nullstellen von $F(R)$ und damit zwei Gebiete in denen $\dot{R}^2 > 0$ ist. Diese entsprechen einem geschlossenen Modell und einem offenen Modell, welches sich jedoch nicht von einer Anfangssingularität heraus entwickelt, sondern von einem Skalenfaktor $R > 0$ ausgehend zunächst auf einen Minimalwert kontrahiert und danach wieder expandiert.

9.5.4 Die deSitter-Lösung

Als letzte der kosmologischen Lösungen betrachten wir eine kosmologische Vakuum-Lösung, die von deSitter (1917) gefunden wurde. Man bekommt dieses Modell, indem man die Friedmann-Gleichung für $\rho = p = k = 0$ löst. In diesem Fall hat man

$$\frac{\dot{R}^2}{R^2} = \frac{1}{3}\Lambda$$

zu lösen. Die Lösung ist einfach zu bekommen und lautet

$$R(t) = R_0 e^{\sqrt{\Lambda/3}t}$$

mit einer beliebigen Integrationskonstanten R_0 . Die Metrik lautet dann

$$g = dt^2 - R_0^2 e^{2\sqrt{\Lambda/3}t} \left(dr^2 + r^2 d\theta^2 + r^2 \sin^2 \theta d\phi^2 \right).$$

Wir führen nun noch neue Koordinaten ein, indem wir

$$x = R_0 r \sin \theta \cos \phi, \quad y = R_0 r \sin \theta \sin \phi, \quad z = R_0 r \cos \theta$$

setzen. Nun lautet die Metrik

$$g = dt^2 - e^{2\sqrt{\Lambda/3}t} (dx^2 + dy^2 + dz^2).$$

Diese Lösung ist die Asymptote aller offenen Modelle mit negativer kosmologischer Konstanten. Dies ist leicht einzusehen, wenn man bedenkt, dass für große R der Term mit ΛR^2 in der Friedmann-Gleichung überwiegt.

9.6 Eine kurze Geschichte des Universums

Zum Abschluss des Kapitels über kosmologische Lösungen wollen wir uns noch kurz anschauen, inwieweit die gefundenen Lösungen uns Informationen liefern über die Entwicklung des Universums vom Beginn an bis heute. Diese Beschreibung ist notgedrungen nur sehr cursorisch und kann den vielen Details nicht Rechnung tragen. Es soll hier nur ein grober Eindruck vermittelt werden. Weitergehende Informationen finden sich in den Büchern von Wald [22] und von Börner [2].

Wir legen unseren Betrachtungen die Standardmodelle mit $\Lambda = 0$ und entweder Staub- oder Strahlungsmaterie zugrunde. Wir hatten gesehen, dass wir für Staub die Relation

$$\rho \sim R^{-3}$$

bekommen, während im Strahlungsuniversum

$$\rho \sim R^{-4}$$

gilt. In unserem Universum gibt es heute sowohl Strahlung als auch Materie. Die heutige Energiedichte der kosmologischen Hintergrundstrahlung wird auf ca. ein Promille der heutigen Energiedichte der normalen Materie geschätzt. Wenn wir die heutigen Werte mit den obigen Relationen in die Vergangenheit extrapolieren, dann waren die Energiedichten bei einem Skalenfaktor von einem Promille des heutigen Wertes von ungefähr gleicher Größenordnung. Noch weiter in der Vergangenheit muss die Energiedichte der Strahlung jedoch gegenüber derjenigen der Materie überwogen haben. Wir sprechen daher von einem **strahlungsdominierten** frühen Universum gegenüber dem **materiedominierten** heutigen Universum. Man wird also annehmen, dass das Strahlungsuniversum eine gute Approximation für das frühe Universum ist und das Universum heute gut durch das materiedominierte Modell wiedergegeben wird.

strahlungsdominiert
materiedominiert

Da die Materie bzw. Strahlung auf immer kleinere Volumina komprimiert wird, wenn man das Universum in die Vergangenheit zurückverfolgt, erwarten wir eine **Temperaturerhöhung** mit abnehmendem Skalenfaktor, die für $R \rightarrow 0$ tatsächlich unendlich groß wird, wie folgende Überlegung zeigt.

Wenn das frühe Universum tatsächlich strahlungsdominiert war, dann erwarten wir, dass für alle Modelle, ungeachtet des Wertes von k , der dominante Term in der Friedmann-Gleichung der Dichteterm ist, so dass in jedem Fall eine Relation

$$R \sim \sqrt{t}, \quad \rho \sim \frac{1}{t^2}$$

resultiert. Eine quantentheoretische Rechnung über die Verteilung von masselosen Teilchen (Photonen) im thermischen Gleichgewicht zeigt, dass für die Energiedichte für Photonen die Temperaturabhängigkeit

$$\rho \sim (kT)^4$$

gilt. Daher gilt für die Temperatur selbst

$$T \sim \rho^{\frac{1}{4}} \sim \frac{1}{R}.$$

Dies legt also die folgende Vorstellung von der Entwicklung unseres Universums nahe. Es begann als ein extrem heißes ($T \rightarrow \infty$), extrem dichtes ($\rho \rightarrow \infty$) Konglomerat von Materie und Strahlung im thermischen Gleichgewicht. Mit der Ausdehnung des Universums nahmen die Distanzen zwischen den Materieteilchen zu, so dass Wechselwirkungen zwischen ihnen weniger häufiger waren. Dadurch kam es zu einer Verletzung des thermischen Gleichgewichts der Materie. Zu dem Zeitpunkt als der Skalenfaktor ein Promille seines heutigen Wertes erreichte, begann die Materiedichte die Strahlungsenergiedichte zu dominieren und die Entwicklung des Universums ähnelt seither immer mehr einem materiedominierten Friedmann-Modell.

Wir beschreiben nun einige wesentliche Stationen in der Entwicklung des Universums aufgrund des sog. heißen Urknallmodells („hot big bang“ model)

1. $t = 1\text{s}$:

Zu diesem Zeitpunkt beträgt die Dichte ca. $5 \times 10^5 \text{g/cm}^3$ und die Temperatur beträgt ca. 10^{10}K . Die Materie besteht jetzt hauptsächlich aus Neutrinos, Photonen, Elektronen, Positronen, Protonen und Neutronen im thermischen Gleichgewicht. Die Temperatur ist klein genug, so dass die Zahl der schwereren Elementarteilchen vernachlässigt werden kann. In diesem Stadium ist die Wechselwirkung mit den Neutrinos so schwach, dass sie vom Rest der Teilchen abkoppeln. Sie werden im weiteren Verlauf der Entwicklung aufgrund der Expansion des Universums rotverschoben, sie verlieren an Energie. Das rotverschobene thermische Spektrum der Neutrinos sollte bis heute eine Temperatur von ca. 2K erreicht haben.

2. $t = 1.5\text{s}$:

Bei weiterer Abkühlung des Universums erreicht die Reaktionsrate des (inversen) β -Zerfalls, bei dem Neutronen in Protonen (und umgekehrt) zerfallen, einen Wert, der wesentlich geringer ist, als die kosmische Expansionsrate. Das bedeutet, dass das thermische Gleichgewicht zwischen Neutronen und Protonen nicht

mehr aufrecht erhalten werden kann. Das Verhältnis zwischen Neutronen und Protonen „friert ein“ auf einem Wert von ca. 1:6. Nach dieser Phase ändert sich dieses Verhältnis nur noch durch den Zerfall der Neutronen.

3. $t = 4s$:

Zu diesem Zeitpunkt beträgt die Dichte ca. $3 \times 10^4 \text{g/cm}^3$ und die Temperatur hat einen Wert von $5 \times 10^9 \text{K} \approx 0.5 \text{MeV}$ erreicht. Dies entspricht ungefähr der Masse des Elektrons oder Positrons. Nun wird es zunehmend schwerer der Paarvernichtung durch Paarerzeugung das Gleichgewicht zu halten. Schließlich überwiegt die Paarvernichtung und bald sind alle Positronen vernichtet. Es bleibt nur noch eine vergleichsweise kleine Zahl von Elektronen zurück.⁵

4. $t = 3m$:

Ab einer Temperatur unter ungefähr 10^9K beginnt die Nukleosynthese, bei der sich die vorhandenen Protonen und Neutronen zu schwereren Kernen, hauptsächlich zu ^4He , verbinden. Nach diesem Prozess hat das erzeugte Helium einen Anteil von ca. 25% an der gesamten Masse der vorhandenen Materie. Die beobachteten Mengen von Helium im Universum scheinen diesen Wert zu unterstützen und es ist schwierig, diesen Wert durch andere Prozesse zu erklären, so dass die Vorhersage von Helium in dieser Größenordnung als ein Erfolg des heißen Urknallmodells gewertet werden muss.

5. $t = 4 \times 10^5 a$ Die Temperatur beträgt nur noch 4000K . Ab diesem Zeitpunkt verbinden sich die übrigen Elektronen und Protonen zu Wasserstoff. Nun gibt es keine freien elektrisch geladenen Teilchen mehr. Als eine Folge entkoppelt die Strahlung von der Materie, weil der Streuquerschnitt von Photonen mit freien geladenen Teilchen wesentlich größer ist als mit neutralem Wasserstoff. Das Universum wurde durchsichtig. Seit dieser Zeit bilden die Photonen die bekannte Hintergrundstrahlung mit einem Spektrum, dessen Temperatur heute 2.7K erreicht hat. Die Existenz dieser Strahlung und die Tatsache, dass sie mit großer Genauigkeit isotrop ist, muss als ein weiterer Erfolg des Modells gewertet werden.

6. $t \approx 10^3 - 10^7 a$:

Das Universum verwandelt sich von strahlungsdominiert zu materiedominiert.

7. $t \approx 1 - 2 \times 10^{10} a$:

Das Universum erreicht seinen heutigen Zustand.

Wie geht es mit dem Universum weiter? Ist es offen oder geschlossen? Diese Frage lässt sich heute noch nicht endgültig beantworten. Wir betrachten noch einmal die Feldglei-

⁵Die Tatsache, dass überhaupt eine der zwei Sorten zurückbleibt, ebenso wie die Tatsache, dass mehr Protonen als Antiprotonen im Universum zu finden sind, dass also keine Materie-Antimaterie-Symmetrie herrscht, ist nicht verstanden. Es kann sein, dass die Wechselwirkungen, die ganz kurz nach dem Urknall wichtig waren, diese Symmetrie verletzen. Die vereinheitlichten Theorien (GUT) scheinen diese Auswirkungen zu haben.

chungen für ein materiedominiertes Weltall

$$3\frac{\dot{R}^2}{R^2} + 3\frac{k}{R^2} - \Lambda = 8\pi\rho$$

$$2\frac{\ddot{R}}{R} + \frac{\dot{R}^2}{R^2} + \frac{k}{R^2} - \Lambda = 0,$$

und schreiben sie mithilfe des Hubble-Parameters $H = \dot{R}/R$ und des (dimensionslosen) Verzögerungsparameters

$$q = -\frac{R\ddot{R}}{\dot{R}^2}$$

in der Form

$$1 = \frac{8\pi\rho}{3H^2} - \frac{k}{H^2R^2} + \frac{\Lambda}{3H^2},$$

$$1 = 2q - \frac{k}{H^2R^2} + \frac{\Lambda}{H^2}.$$

Im Fall $\Lambda = 0$, den wir hier besprechen, folgt $q = \frac{4\pi\rho}{3H^2}$. Das Weltall ist geschlossen, also $k \geq 0$ genau dann, wenn

$$\rho > \rho_c = \frac{3H^2}{8\pi} \iff q > \frac{1}{2}$$

gilt. Die kritische Dichte ρ_c wird also durch den Hubble-Parameter fixiert. Man definiert üblicherweise das Verhältnis von Dichte zu kritischer Dichte als

$$\Omega := \rho/\rho_c.$$

Die Entscheidung, ob das Universum offen oder geschlossen ist, muss durch Beobachtung festgestellt werden. Es gibt zur Zeit vier unabhängige Beobachtungsmethoden, die diese Frage in der Zukunft vielleicht beantworten können, wenn die Messungen genügende Genauigkeit erreicht haben werden. Dies sind:

- Die Abweichung von der linearen Beziehung zwischen Rotverschiebung und Helligkeitsabstand. Diese werden im wesentlichen verursacht durch die Dynamik des Universums und geben so einen Hinweis auf den Verzögerungsparameter q . Ist $q \leq 1/2$, dann ist das Universum offen, sonst geschlossen. Diese Abweichungen zeigen sich jedoch erst für Objekte, die sehr weit von uns entfernt sind, die wir also in einem sehr frühen Stadium sehen. Es ist nicht geklärt, inwieweit diese sehr frühen Objekte als Standardkerzen dienen können.
- Die Dichte der Materie im Universum. Der Parameter Ω lässt sich wesentlich genauer bestimmen, als ρ oder H für sich alleine, die beide durch die Ungenauigkeit der Abstandsmessung beeinträchtigt sind. Diese heben sich bei der Bestimmung von Ω gerade heraus. Die momentanen Messungen ergeben einen Wert von $\Omega \approx 0.04$. Diese enthalten jedoch nur den sichtbaren Teil der Materie. Es gibt Hinweise, dass unser Universum auch **dunkle Materie** enthält. Dies ergibt

dunkle Materie

sich aus Messungen der Geschwindigkeiten von Sternen in Galaxien. Diese sind für Sterne fern vom galaktischen Zentrum wesentlich höher als sie sein dürften, wenn die Galaxis nur aus den sichtbaren Sternen bestünde.

- Das Alter des Universums. Für das materiedominierte Universum mit $k = 0$ ist das Alter des Universums mit dem Hubble-Parameter über

$$T = \frac{2}{3H}$$

verknüpft. Für das geschlossene Universum liegt der Wert darüber, für das offene Universum darunter. Eine Abschätzung für das Mindestalter des Universums ergäbe also auch eine Aussage über das zukünftige Verhalten des Universums. Leider sind die entsprechenden Messungen nicht geeignet, um solche Aussagen zu treffen.

Neben diesen Messungen der verschiedenen individuellen kosmologischen Parameter gibt es neuerdings auch Resultate, die auf der Vermessung der Anisotropien der kosmischen Hintergrundstrahlung beruhen. Diese ist zu einem hohen Grade isotrop und rechtfertigt daher die Annahme eines isotropen und homogenen Universums. Dennoch ist unser Universum auf kleineren Skalen nicht isotrop und diese Anisotropien hinterlassen Spuren in der Hintergrundstrahlung, die ja durch dieses leicht anisotrope Universum zu uns reist. Durch genaue Vermessung dieser Anisotropien in der Hintergrundstrahlung und durch die Berechnung der verschiedenen Effekte, die die Strahlung bei ihrer Propagation durch das anisotrope Universum beeinflussen, kann man unter Zugrundelegung eines bestimmten kosmologischen Modells die Parameter dieses Modells den Beobachtungen anpassen und damit 'messen'. Die momentanen Werte, die auf Messungen des Satelliten **WMAP** beruhen sind:

WMAP

1. Das Universum ist 13.7 Milliarden Jahre alt.
2. Das Universum hat einen Durchmesser von mindestens 78 Milliarden Lichtjahren.
3. Das Universum ist flach, $k = 0$.
4. Das Universum besteht aus 4% sichtbarer Materie, 23% dunkler Materie sowie 73% 'dunkler Energie'.
5. Die Hubble-Konstante beträgt $71 \pm 4 \text{ km/s/Mpc}$.

Die angesprochene dunkle Energie ist ein Phänomen, das im Moment nicht erklärt werden kann. Eine Erklärung ist, diese Form von Energie der kosmologischen Konstanten zuzuschreiben. Andere postulieren eine neue Materieform namens „Quintessenz“. Was von diesen Versuchen zu halten ist, ist unklar. Fest steht, dass die Messungen von WMAP zu den genauesten gehört, die bisher durchgeführt wurden. Inwieweit die Anpassung an die Modellparameter Bestand hat, inwieweit diese durch gewisse Vorurteile beeinflusst sind, muss man noch sehen.

10 Eigenschaften von Schwarzen Löchern

In diesem Kapitel sollen noch einmal die Schwarzen Löcher behandelt werden. Wir werden uns insbesondere ihre globalen Eigenschaften und ihre Kausalitätsstruktur anschauen. Außerdem werden verschiedene Sorten von Schwarzen Löchern diskutiert.

10.1 Die Kruskal-Raumzeit

Wir kommen noch einmal auf die Schwarzschild-Lösung zurück. In Kapitel 7 hatten wir gesehen, dass die Schwarzschild-Metrik die Form

$$ds^2 = \left(1 - \frac{2m}{r}\right) dt^2 - \frac{1}{\left(1 - \frac{2m}{r}\right)} dr^2 - r^2 (d\theta^2 + \sin^2 \theta d\phi^2) \quad (10.1.1)$$

hat. Wie schon dort festgestellt wurde, ist diese Form der Metrik *singulär* bei $r = 2m$ und bei $r = 0$. Die Metrik in dieser Form gilt also nur im Bereich $r > 2m$ oder im Bereich $r < 2m$, aber nicht in beiden gemeinsam, weil die Komponente g_{rr} bei $r = 2m$ nicht definiert ist.

Während bei $r = 0$ tatsächlich eine echte physikalische Singularität vorliegt, handelt es sich bei der „Problemstelle“ $r = 2m$ jedoch um eine Koordinatensingularität. Dies hatten wir festgestellt, weil es möglich war, andere Koordinaten zu finden, so dass in diesen neuen Koordinaten bei der entsprechenden Stelle keine Singularität mehr vorhanden war.

Die entsprechenden Koordinaten sind die *avancierten* bzw. *retardierten* **Eddington-Finkelstein Koordinaten**. Man erhält die avancierten (+) bzw. retardierten (−) Koordinaten durch Einführung der neuen Koordinate

Eddington-Finkelstein Koordinaten

$$w = t \pm \left(r + 2m \log\left(\frac{r}{2m} - 1\right)\right), \quad \text{für } r > 2m$$

anstelle von t . In diesem neuen Koordinatensystem lautet die Metrik

$$ds^2 = \left(1 - \frac{2m}{r}\right) dw^2 \mp 2dw dr - r^2 (d\theta^2 + \sin^2 \theta d\phi^2).$$

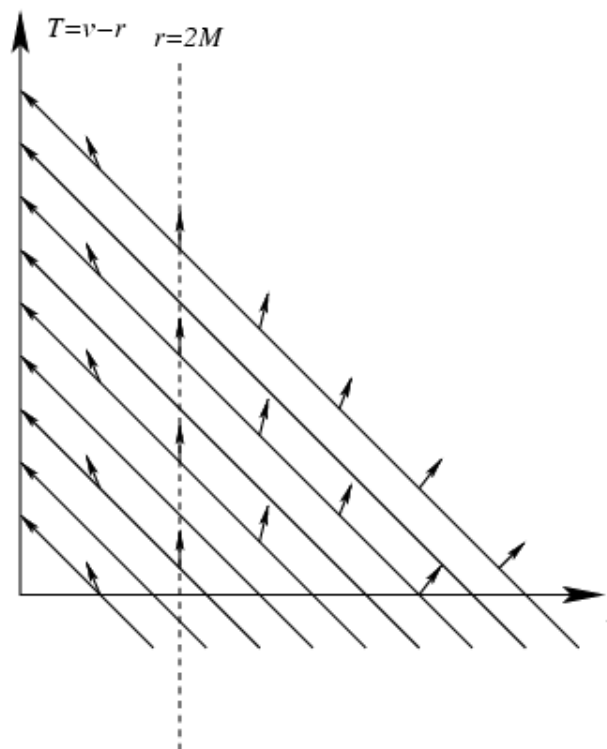


Abbildung 10.1: Die Raum-Zeit in avancierten Eddington-Finkelstein Koordinaten

Offensichtlich ist an der Stelle $r = 2m$ nun keine Singularität mehr zu entdecken. Das Koordinatensystem (w, r, θ, ϕ) , das ursprünglich nur für $r > 2m$ definiert war, kann man nun sogar für Werte $r < 2m$ ausdehnen. Die Metrik behält dabei ihre Form und erfüllt natürlich weiterhin die Einstein-Gleichungen. Wir haben somit den Gültigkeitsbereich der Schwarzschild-Metrik *fortgesetzt*. Diese Form der Metrik ist gültig für alle Werte $r > 0$. Kann man diesen Prozess fortsetzen? D.h. ist es möglich, eine weitere Koordinatentransformation zu finden, die uns vielleicht sogar weitere Gebiete erschließt? Gibt es ein Kriterium, welches uns sagt, ob bzw. wann dieser Prozess abgeschlossen ist? Wir werden versuchen, in diesem Kapitel Antworten auf diese Fragen zu finden.

Wir setzen im Folgenden $w = u$ für die retardierten und $w = v$ für die avancierten Eddington-Finkelstein Koordinaten. In Abb. 7.2, die wir hier noch einmal zeigen, ist die Raum-Zeit in avancierten Eddington-Finkelstein Koordinaten gezeigt. Das Wesentliche an diesem Diagramm ist die Tatsache, dass die lichtartigen radial einlaufenden Geodäten als Geraden dargestellt sind. Dies folgt aus der einfachen Überlegung, dass radiale Kurven einen Tangentialvektor u^a mit der Darstellung

$$u = a\partial_v + b\partial_r$$

besitzen müssen. Lichtartige Kurven haben einen lichtartigen Tangentialvektor, es gilt

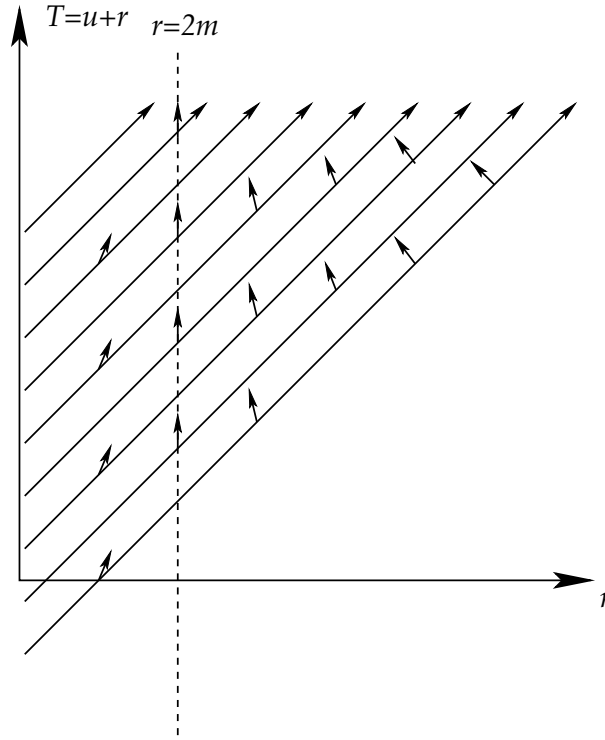


Abbildung 10.2: Die Raum-Zeit in retardierten Eddington-Finkelstein Koordinaten

also

$$0 = g_{ab}u^a u^b = a^2 \left(1 - \frac{2m}{r}\right) - 2ab,$$

also $a = 0$ oder $b = a \left(1 - \frac{2m}{r}\right)$. Der Fall $a = 0$ führt auf die einlaufenden Geodäten mit

$$u^a \sim \partial_r^a.$$

Für diese Klasse von Kurven bleibt v konstant. Dies sind die Geraden in Abb. 10.1. Die zweite Klasse von Geodäten besitzt einen Tangentialvektor

$$u^a \sim \partial_v^a + \frac{1}{2} \left(1 - \frac{2m}{r}\right) \partial_r^a.$$

Dieser ist in der Abb. 10.1 ebenfalls angegeben. Dabei ist festzustellen, dass der Tangentialvektor im Bereich $r > 2m$ nach außen zeigt, während er innerhalb von $r = 2m$ ebenfalls nach innen zeigt.

Das entsprechende Diagramm für retardierte Eddington-Finkelstein Koordinaten ist in Abb. 10.2 gezeigt. Hier sind es die radial *auslaufenden* Geodäten, die als Geraden dargestellt sind. Wir stellen fest, dass bei Einführung der neuen Koordinate $W = u$ bzw.

$w = v$ die eine oder andere Klasse von Geodäten eine einfache Beschreibung durch die Koordinaten zulässt.

Wir bestimmen die Gleichung der radialen Geodäten in retardierten Eddington-Finkelstein Koordinaten. Es ergibt sich

$$\frac{d}{d\lambda} \left[\dot{u} \left(1 - \frac{2m}{r} \right) + \dot{r} \right] = 0, \quad (10.1.2)$$

$$\ddot{u} = -\frac{m}{r^2} \dot{u}^2. \quad (10.1.3)$$

Wir haben außerdem noch die Bedingung, dass der Tangentialvektor $\dot{u}\partial_u + \dot{r}\partial_r$ lichtartig sein muss, und diese lautet

$$\dot{u}^2 \left(1 - \frac{2m}{r} \right) + 2\dot{r}\dot{u} = 0.$$

Die erste Klasse von Geodäten ist durch $\dot{u} = 0$ gegeben und für diese folgt aus (10.1.2) $\ddot{r} = 0$, also $r(\lambda) = a\lambda + b$. Zukunftsgerichtete Kurven dieser Klasse mit Anfangspunkt (u_0, r_0) sind *auslaufend*, also $\dot{r} > 0$ und wir erhalten (mit willkürlicher Wahl $a = 1$)

$$r(\lambda) = r_0 + \lambda, \quad u(\lambda) = u_0.$$

Die zweite Klasse von Geodäten erfüllt die Gleichung

$$\dot{u} = -\frac{2\dot{r}}{1 - \frac{2m}{r}},$$

die sich leicht integrieren lässt zu

$$u(r) - u_0 = -2 \left(r - r_0 + 2m \log \left(\frac{r - 2m}{r_0 - 2m} \right) \right).$$

Dies liefert uns also u in Abhängigkeit von r . Andererseits ist aber mithilfe von (10.1.2) auch hier $\ddot{r} = 0$. Hier sind zukunftsgerichtete Kurven (für $r > 2m$) *einlaufend*, also $\dot{r} < 0$ und wir bekommen die Darstellung

$$r(\lambda) = r_0 - \lambda, \quad u(\lambda) = 2 \left(\lambda - 2m \log \left(\frac{r_0 - 2m - \lambda}{r_0 - 2m} \right) \right).$$

Verfolgen wir die beiden Klassen von Geodäten in Abhängigkeit von ihrem Parameter. Für die auslaufenden Geodäten ergeben sich keine Probleme. Sie kommen von $r = 0$ her, welches sie bei einem endlichen Parameterwert $\lambda = -r_0$ 'verlassen' und laufen mit wachsendem Parameter ins Unendliche, welches sie für $\lambda \rightarrow \infty$ auch erreichen.

Bei den einlaufenden Geodäten hingegen stoßen wir auf Probleme. Wir sehen, dass für $\lambda = r_0 - 2m$, also bei $r = 2m$ die Darstellung dieser Kurven nicht mehr definiert ist, weil dort der Logarithmus divergiert. Es gilt also $v \rightarrow \infty$, obwohl λ einen endlichen Wert erreicht und obwohl die Metrik dort nicht singulär ist. Dies deutet auf eine schlechte Wahl

der Koordinaten hin. Hätten wir die Geodäten in avancierten Eddington-Finkelstein Koordinaten betrachtet, dann hätten wir ähnliche Ergebnisse gefunden. Eine Sorte der Geodäten ist gutartig, die andere erreicht unendliche Koordinatenwerte bei endlichem affinem Parameter. Wie können wir beide Klassen von Geodäten 'bezähmen'?

Es liegt nahe, nicht nur *entweder* avancierte *oder* retardierte Koordinaten einzuführen, sondern *beide*. Wir machen also die Koordinatentransformation für $r > 2m$

$$\begin{aligned}u &= t - \left(r + 2m \log\left(\frac{r}{2m} - 1\right) \right) \\v &= t + \left(r + 2m \log\left(\frac{r}{2m} - 1\right) \right).\end{aligned}$$

Die Umkehrung dieser Transformation, die t und r in Abhängigkeit von u und v gibt, läßt sich nicht vollständig explizit hinschreiben. Vielmehr gilt

$$\begin{aligned}\frac{1}{2}(v + u) &= t, \\ \frac{1}{2}(v - u) &= r + 2m \log\left(\frac{r}{2m} - 1\right).\end{aligned}$$

Die zweite Relation kann leider nicht explizit nach r aufgelöst werden. Trotzdem können wir $r = r(u, v)$ als Funktion von u und v betrachten.

Mit dieser Transformation wird die Metrik

$$ds^2 = \left(1 - \frac{2m}{r(u, v)}\right) du dv - r(u, v)^2 (d\theta^2 + \sin^2 \theta d\phi^2).$$

Übung 10.1: Diskutieren Sie das Verhalten von radialen Lichtstrahlen in diesem Koordinatensystem. Wie verlaufen die Geodäten in einem (T, R) -Diagramm, wo $T = \frac{1}{2}(u + v)$ und $R = \frac{1}{2}(v - u)$ gilt? Wie sieht die Parametrisierung $(u(\lambda), v(\lambda))$ aus?

Die einzige Koordinatentransformation, die diese Form der Metrik invariant lassen, sind

$$u = u(U), \quad v = v(V).$$

Unter dieser Transformation ändert sich der (u, v) -Teil der Metrik zu

$$\left(1 - \frac{2m}{r(u, v)}\right) du dv = \left(1 - \frac{2m}{r(U, V)}\right) \frac{du}{dU} \frac{dv}{dV} dU dV.$$

Es ändert sich also nur der Vorfaktor. Man kann diese Freiheit in der Wahl der Null-Koordinaten benutzen, um den Vorfaktor in eine geeignete Form zu bringen. Kruskal wählte als neue Null-Koordinaten die Funktionen

$$\begin{aligned}U &= -e^{-\frac{u}{4m}} \iff u = -4m \log(-U) \\ V &= e^{\frac{v}{4m}} \iff v = 4m \log(V).\end{aligned}$$

Mit diesen Funktionen gilt

$$dU = -U \frac{du}{4m}, \quad dV = V \frac{dv}{4m}$$

so dass wir nach etwas Rechnung die Metrik in der Form

$$ds^2 = 16m^2 \left(\frac{2m}{r} \right) e^{-r/2m} dU dV - r^2 (d\theta^2 + \sin^2 \theta d\phi^2) \quad (10.1.4)$$

erhalten. Mit einer letzten Koordinatentransformation

$$U = T - X, \quad V = T + X$$

bringen wir die Metrik schließlich auf die Form

$$ds^2 = 16m^2 \left(\frac{2m}{r} \right) e^{-r/2m} (dT^2 - dX^2) - r^2 (d\theta^2 + \sin^2 \theta d\phi^2). \quad (10.1.5)$$

Die Beziehung zwischen den ursprünglichen Koordinaten (t, r) und den neuen Koordinaten (T, X) ergibt sich aus den Gleichungen

$$T^2 - X^2 = -e^{r/2m} \left(\frac{r}{2m} - 1 \right), \quad (10.1.6)$$

$$T/X = \tanh\left(\frac{t}{4m}\right). \quad (10.1.7)$$

Während man diese Relationen nach t auflösen kann, ist dies für r nicht möglich. Dieses ist immer noch als implizite Funktion von T und X zu sehen, definiert durch die erste dieser Gleichungen.

Offensichtlich ist diese Form der Metrik ebenfalls regulär für alle Werte von T und X , für die $r > 0$ gilt. Betrachten wir nun die radialen Lichtstrahlen, dann stellen wir schnell fest, dass diese durch $U = \text{const}$ oder $V = \text{const}$ charakterisiert sind. In einem (T, X) -Diagramm (siehe Abb. 10.3) laufen sie daher auf Geraden mit Steigung $+1$ bzw. -1 , also unter 45° . Die Linien $t = \text{const}$ werden ebenfalls durch Geraden dargestellt, jedoch mit der Steigung $\tanh(\frac{t}{4m})$, was betragsmäßig immer kleiner als 1 ist. Diese Geraden laufen alle durch den Ursprung $T = 0 = X$. Die Kurven auf denen r konstant ist, sind Hyperbeln mit jeweils zwei Ästen und zwar zeitartige Hyperbeln für $r > 2m$ und raumartige Hyperbeln für $r < 2m$. Die Kurve $r = 0$ besteht demnach aus den Punkten mit

$$T^2 - X^2 = 1,$$

also zwei Hyperbelästen, zwischen denen sich die Raumzeit befindet. Die Punkte auf der Hyperbel selbst gehören nicht mehr Raumzeit, denn die Metrik ist dort nicht mehr definiert. Wir bekommen also *zwei* singuläre Gebiete, die $r = 0$ entsprechen.

Dies ist ein allgemeines Phänomen. Für jeden Wert von r gibt es nicht nur eine Hyperfläche mit konstantem r , sondern zwei. Es ist als ob sich die ursprüngliche Raumzeit verdoppelt hätte.

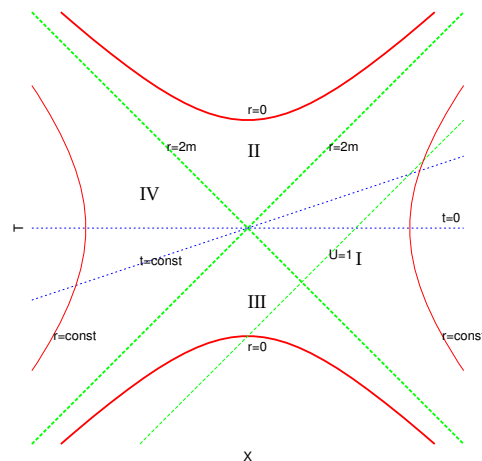


Abbildung 10.3: Die Kruskal-Raumzeit

Das Raum-Zeit Diagramm der Kruskal-Raumzeit besteht aus vier Gebieten, die durch die Geraden $U = 0$ bzw. $V = 0$ voneinander getrennt sind. Die ursprüngliche Koordinatentransformation war beschränkt auf den Bereich $r > 2m$, das entspricht $U < 0$ und $V > 0$. Dies ist genau der Bereich I. Es ist genau dieser Bereich, den uns die ursprünglichen Koordinaten zu sehen gestatten. Dieses Koordinatensystem entartet an den Grenzen, wo $r = 2m$ wird.

schwarzes Loch

Das Gebiet II enthält eine der beiden singulären Hyperbeln $r = 0$. Es ist leicht einzusehen, dass jede zeit- oder lichtartige Kurve, die aus I oder III kommend in II eintritt, dieses Gebiet nicht mehr verlassen kann. Diese Kurve wird unweigerlich in der Singularität landen. Es handelt sich also um eine zukünftige Singularität. Da sie raumartig ist, kann sie von den darauf zufallenden Beobachtern nicht bemerkt werden. Sie erkennen sie erst in dem Moment, wenn sie auf sie treffen. Dieses Gebiet ist hinter dem Horizont bei $r = 2m$ versteckt. Es handelt sich um ein **schwarzes Loch**.

weißes Loch

Gebiet III enthält ebenfalls eine Singularität mit $r = 0$. Im Gegensatz zu II ist diese Singularität in der Vergangenheit. Das bedeutet, dass jede zeit- oder lichtartige Kurve die aus III in I oder IV übergeht, aus der Singularität bei $r = 0$ gekommen sein muss. Es handelt sich daher um ein **weißes Loch**.

asymptotische
Region

Gebiet IV schließlich ist in Analogie zu I zu sehen. Es ist ein Gebiet, in dem der Radius unbegrenzt wachsen kann. Es handelt sich in den beiden Fällen um eine **asymptotische Region**.

10.2 Die Einstein-Rosen Brücke

Das Kruskal-Diagramm Abb. 10.3 sieht einfach aus. Wir dürfen aber nicht vergessen, dass jeder Punkt (T, X) in diesem Bild eine ganze Kugeloberfläche S^2 repräsentiert, die einen Radius r besitzt, der durch Gleichung (10.1.6) bestimmt ist. Nur die Punkte auf der Singularität $r = 0$ sind einzelne Punkte, keine Kugeloberflächen.

Dies hat zur Folge, dass die Struktur der Kruskal-Raumzeit doch etwas komplizierter ist, als es zuerst den Anschein hat. Betrachten wir zuerst die Hyperfläche $t = 0$. Dies ist im Kruskal-Diagramm die Hyperfläche $T = 0$. Die Metrik dieser Fläche bekommt man, indem man in der Kruskal-Metrik (10.1.5) $dT = 0$ setzt und dann das Vorzeichen ändert, also

$$h = 16m^2 \left(\frac{2m}{r} \right) e^{-r/2m} dX^2 + r^2 (d\theta^2 + \sin^2 \theta d\phi^2).$$

Was als erstes auffällt, ist das Verhalten von r wenn X zwischen ∞ und $-\infty$ variiert. Für große X ist auch r groß und schrumpft mit abnehmendem X . Für $X = 0$ ist aber $r = 2m \neq 0$. Für weiter abnehmende X wächst r wieder an. Wir finden also eine **Minimalfläche**, eine Fläche minimalen Flächeninhalts.

Minimalfläche

Um die Struktur dieser Hyperfläche, die ja ein drei-dimensionaler kugelsymmetrischer Raum ist, noch etwas genauer zu sehen, beschränken wir uns auf die Äquatorebene, setzen also $\theta = \pi/2$. Diese Ebene besitzt die Metrik

$$16m^2 \left(\frac{2m}{r} \right) e^{-r/2m} dX^2 + r^2 d\phi^2.$$

Aus allen Kugeloberflächen $r = \text{const}$, insbesondere aus der Minimalfläche ist jetzt ein Kreis mit Radius r geworden. Diese Kreise stecken alle ineinander wie in der euklidischen Ebene, sie folgen aber aufgrund des Vorfaktors vor dX^2 unterschiedlich 'dicht' aufeinander. Um zu sehen, welche Geometrie sich in dieser Aufeinanderfolge von Kreisen verbirgt, ziehen wir sie auseinander wie einen japanischen Lampion. Mathematisch gesprochen, versuchen wir eine axialsymmetrische Einbettung der Äquatorebene in den drei-dimensionalen euklidischen Raum zu finden.

Die Metrik des drei-dimensionalen euklidischen Raums lautet in Zylinderkoordinaten

$$d\rho^2 + dz^2 + \rho^2 d\phi^2.$$

Eine axialsymmetrische Fläche wird beschrieben durch eine Funktion $\rho = h(z)$, die zu jeder Höhe z den Abstand von der Rotationsachse zur Fläche angibt. Damit berechnet man die Metrik dieser Fläche zu

$$(1 + h'(z)^2) dz^2 + h(z)^2 d\phi^2.$$

Die Bedingung, dass diese Metrik mit der Metrik der Äquatorialebene übereinstimmt, dass also

$$(1 + h'(z)^2) dz^2 + h(z)^2 d\phi^2 = 16m^2 \left(\frac{2m}{r} \right) e^{-r/2m} dX^2 + r^2 d\phi^2$$

gilt, liefert uns zunächst die Gleichung

$$h(z) = r.$$

Wir versuchen nun z als Funktion von r zu bestimmen. Der Vergleich der anderen Koeffizienten liefert

$$(1 + h'(z)^2) \left(\frac{dz}{dr} \right)^2 = 16m^2 \left(\frac{2m}{r} \right) e^{-r/2m} \left(\frac{dX}{dr} \right)^2.$$

Wir setzen für die weiteren Überlegungen $2m = 1$ (bzw. wir messen r und z in Einheiten von $2m$) und bekommen dann mithilfe von (10.1.6) die Gleichung

$$\frac{dz}{dr} = \pm \sqrt{\frac{1 + T^2 e^{-r}}{r - (1 + T^2 e^{-r})}},$$

welche sich für $T = 0$ explizit lösen lässt. Man bekommt $\rho = r(z)$ mit

$$r(z) = 1 + \frac{z^2}{4}.$$

Einstein-Rosen-
Brücke
Wormloch

Diese Gleichung beschreibt eine Rotationsfläche, die durch Rotation einer Parabel um die z -Achse entsteht, siehe Abb. 10.4. Dieses Bild zeigt die **Einstein-Rosen-Brücke**, die manchmal auch als das Schwarzschild-**Wormloch** bezeichnet wird. Wir sehen zwei asymptotisch flache Gebiete, die durch die Brücke miteinander verbunden sind. Das untere Gebiet entspricht dem Teil der $T = 0$ Hyperfläche, der im Gebiet I des Kruskal-Diagramms liegt, während der obere Teil dem Gebiet IV entspricht. Wenn wir uns also auf dem Kruskal-Diagramm auf der Linie $T = 0$ von positiven zu negativen X -Werten bewegen, dann bewegen wir uns in dem Einbettungsdiagramm von unten nach oben. Der Radius der durchquerten Kugelflächen wird dabei immer kleiner, erreicht bei $X = 0$ sein Minimum $r = 2m$ und wächst dann wieder an.

Betrachten wir die entsprechenden Rotationsfiguren für $T \neq 0$, so bekommen wir die folgenden Diagramme.

Wir sehen also das folgende Bild. Die Kruskal-Raumzeit beschreibt zwei asymptotisch flache Universen beschrieben durch $X < 0$ und $X > 0$, die anfänglich voneinander getrennt sind und jedes eine Singularität bei $r = 0$ besitzt. Im Verlauf der Zeit T 'näheren' sich die beiden Singularitäten aneinander an, und vereinigen sich schließlich bei $T = -1$. In diesem Moment entsteht eine Verbindung zwischen den beiden Universen, die immer größer wird und bei $T = 0$ schließlich ein Maximum erreicht. Danach schließt sich diese Verbindung wieder und die Singularitäten trennen sich voneinander.

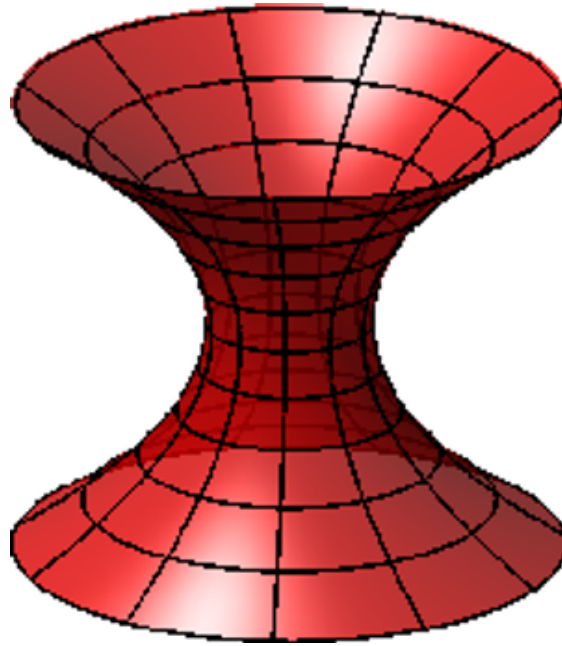


Abbildung 10.4: Einbettungsdiagramm der Äquatorialebene für $T = 0$

Kann man durch diese Verbindung reisen? Zur Beantwortung dieser Frage betrachten wir noch einmal das Kruskal-Diagramm in Abb. 10.3 und beachten, dass in diesem Diagramm die Lichtstrahlen auf Geraden mit Steigung $\pm 45^\circ$ laufen. 'Reisen' bedeutet wie immer, sich auf einer zeitartigen Kurve zu bewegen, deren Steigung folglich immer zwischen -45° und $+45^\circ$ liegen muss. Unter diesen Vorgaben ist es leicht einzusehen, dass es unmöglich ist, eine zeitartige Kurve zu finden, die vom Gebiet I ins Gebiet IV führt. Es ist also *nicht* möglich, von einem Universum ins andere durch die Einstein-Rosen-Brücke zu gelangen. Die Verbindung öffnet und schließt sich so schnell, dass es kein Beobachter schafft, rechtzeitig hindurch zu kommen.

Man kann sich nun fragen, ob es möglich ist, Lösungen der Einstein-Gleichungen zu finden, in denen es ein ähnliches Wurmloch gibt, welches aber so lange offen bleibt, dass man tatsächlich hindurchreisen kann? Diese Frage ist schwierig zu beantworten, wenn man nur physikalisch sinnvolle Lösungen betrachten will. Es ist einfach, eine Metrik hinzuschreiben, die solche Eigenschaften hat. Wenn man nun jedoch den Einstein-Tensor dieser Metrik berechnet und daraus via Einstein-Gleichungen den Energie-Impuls-Tensor bestimmt, der diese Metrik erzeugt, dann findet man, dass dieser nicht physikalisch sinnvoll ist. Dies äußert sich z.B. darin, dass man negative Energiedichte und Drücke findet. Es handelt sich also um Material, das man zumindest bis heute nicht so ohne weiteres im Labor herstellen kann. Trotz dieser eher ins Science-Fiction Genre abgleitenden Bemerkung sei hier die einfachste Metrik angegeben, die ein statisches Wurmloch beschreibt:

$$g = dt^2 - dx^2 - (x^2 + a^2)(d\theta^2 + \sin^2 \theta d\phi^2)$$

Übung 10.2: Bestimmen Sie den Energie-Impuls-Tensor, der diese Metrik erzeugt. Diskutieren Sie den Verlauf von radialen Lichtstrahlen in dieser Metrik. Wie sieht das Einbettungsdiagramm der Äquatorialebene einer $t = \text{const}$ Hyperfläche aus?

10.3 Die Kerr-Lösung

Kerr-Lösung

Zum Abschluß des Kapitels über Schwarze Löcher kommen wir nun zu einer exakten Lösung der Einstein-Gleichungen, die wie keine andere Einfluß auf die Entwicklung der Gravitationsphysik genommen hat. Die **Kerr-Lösung** ist eine Vakuum-Lösung, die die Schwarzschild-Lösung verallgemeinert. Sie beschreibt das Gravitationsfeld einer bestimmten Konfiguration einer rotierenden Masse, so wie die Schwarzschild-Lösung eine sphärisch symmetrische Massenverteilung beschreibt.

Im Gegensatz zur Schwarzschild-Lösung gibt es jedoch keine reguläre innere Lösung mit idealer Flüssigkeit, die an die Kerr-Lösung angeschlossen werden könnte. In diesem Sinne kann man nicht sagen, dass diese Lösung von einer Materieverteilung erzeugt würde. Die wahre Bedeutung der Kerr-Lösung liegt darin, dass sie den *Endzustand* eines Schwarzen Lochs charakterisiert, das durch den Gravitationskollaps einer rotierenden Masse z.B. eines Neutronensterns entstanden. Nach einer gewissen Zeit wird dieses Schwarze Loch alle „Unregelmäßigkeiten“, die von dem kollabierten Körper herrühren in Form von Strahlung und anderen Prozessen ‘abgeschüttelt’ haben und zur Ruhe kommen. Es wird sich also einem stationären Gleichgewichtszustand annähern. Dieser wird durch die Kerr-Lösung beschrieben. Dieses Szenario ist zwar nicht lückenlos bewiesen, viele Untersuchungen deuten jedoch sehr stark darauf hin.

In diesem Kapitel werden diese Lösung etwas genauer untersuchen und einige ihrer Eigenschaften erläutern. Der Weg von den Einstein-Gleichungen zur Kerr-Lösung ist nicht gerade einfach. Wir werden uns daher darauf beschränken, die Metrik anzugeben, ohne die Gleichungen explizit zu lösen. Als Einstieg wollen wir uns aber einen ‘Trick’ anschauen, der aus der Schwarzschild-Lösung die Kerr-Lösung generiert.

10.3.1 Null-Tetraden

Dieser Trick beruht auf der Verwendung von Null-Tetraden. Dies ist eine spezielle Basis im Tangentialraum eines Punktes $P \in \mathcal{M}$, die wie folgt konstruiert werden kann. Wir beginnen mit einer beliebigen Orthonormalbasis

$$(\mathbf{e}_0, \mathbf{e}_1, \mathbf{e}_2, \mathbf{e}_3),$$

also vier Einheitsvektoren mit Skalarprodukt

$$g_{ij} = g(\mathbf{e}_i, \mathbf{e}_j) = \eta_{ij} = \text{diag}(+1, -1, -1, -1).$$

Nun definieren wir zwei lichtartige Vektoren

$$\mathbf{l} = \frac{1}{\sqrt{2}}(\mathbf{e}_0 + \mathbf{e}_1), \quad \mathbf{n} = \frac{1}{\sqrt{2}}(\mathbf{e}_0 - \mathbf{e}_1)$$

Die übrigen beiden raumartigen Vektoren kombinieren wir ebenfalls zu lichtartigen Vektoren. Dies kann man mit reellen Koeffizienten nicht erreichen,

Übung 10.3: Warum nicht?

daher bilden wir formale komplexe Linearkombinationen

$$\mathbf{m} = \frac{1}{\sqrt{2}}(\mathbf{e}_2 + i\mathbf{e}_3), \quad \bar{\mathbf{m}} = \frac{1}{\sqrt{2}}(\mathbf{e}_2 - i\mathbf{e}_3).$$

Nun ist tatsächlich

$$g(\mathbf{m}, \mathbf{m}) = 0, \quad g(\bar{\mathbf{m}}, \bar{\mathbf{m}}) = 0, \quad g(\bar{\mathbf{m}}, \mathbf{m}) = -1.$$

Insgesamt ist die Matrix der Skalarprodukte der Basis $(\mathbf{l}, \mathbf{n}, \mathbf{m}, \bar{\mathbf{m}})$

$$g_{ij} = \begin{pmatrix} 0 & 1 & 0 & 0 \\ 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & -1 \\ 0 & 0 & -1 & 0 \end{pmatrix}$$

Eine andere Schreibweise des gleichen Sachverhalts bekommen wir mit der Indexschreibweise

$$l_a l^a = n_a n^a = m_a m^a = \bar{m}_a \bar{m}^a = 0, \quad l_a n^a = 1 = -m_a \bar{m}^a. \quad (10.3.1)$$

Eine solche Basis nennt man **Null-Tetrad**. Man kann mithilfe einer Null-Tetrad die Metrik darstellen als Null-Tetrad

$$g_{ab} = l_a n_b + n_a l_b - m_a \bar{m}_b - \bar{m}_a m_b. \quad (10.3.2)$$

Übung 10.4: Zeigen Sie die Gültigkeit dieser Darstellung, indem sie die Skalarprodukte der Basisvektoren berechnen. Warum reicht dieser Nachweis aus?

Ganz entsprechend läßt sich die inverse Metrik darstellen

$$g^{ab} = l^a n^b + n^a l^b - m^a \bar{m}^b - \bar{m}^a m^b.$$

Wenn man also eine (Null)-Tetrad kennt, dann kennt man auch die Metrik. Wir wollen uns dies hier zunutze machen und die Schwarzschild-Metrik in avancierten Eddington-Finkelstein Koordinaten betrachten

$$ds^2 = \left(1 - \frac{2m}{r}\right) dv^2 - 2dvdr - r^2 (d\theta^2 + \sin^2 \theta d\phi^2).$$

Es ist schnell einzusehen, dass folgende Vektoren eine Null-Tetrad bilden

$$\begin{aligned} \mathbf{l} &= \partial_r, & \mathbf{m} &= \frac{1}{\sqrt{2}r} \left(\partial_\theta + \frac{i}{\sin \theta} \partial_\phi \right), \\ \mathbf{n} &= -\partial_v - \frac{1}{2} \left(1 - \frac{2m}{r} \right) \partial_r, & \bar{\mathbf{m}} &= \frac{1}{\sqrt{2}r} \left(\partial_\theta - \frac{i}{\sin \theta} \partial_\phi \right). \end{aligned}$$

Der 'Trick' beginnt damit, dass wir r erlauben komplexe Werte anzunehmen. Wir schreiben die Tetrad nun in der Form

$$\begin{aligned} \mathbf{l} &= \partial_r, & \mathbf{m} &= \frac{1}{\sqrt{2}\bar{r}} \left(\partial_\theta + \frac{i}{\sin \theta} \partial_\phi \right), \\ \mathbf{n} &= -\partial_v - \frac{1}{2} \left(1 - \frac{m}{r} - \frac{m}{\bar{r}} \right) \partial_r, & \bar{\mathbf{m}} &= \frac{1}{\sqrt{2}r} \left(\partial_\theta - \frac{i}{\sin \theta} \partial_\phi \right). \end{aligned}$$

Als nächstes machen wir eine Koordinatentransformation

$$v \mapsto v' = v + ia \cos \theta, \quad r \mapsto r' = r + ia \cos \theta, \quad \theta \mapsto \theta' = \theta, \quad \phi \mapsto \phi' = \phi.$$

Dann ist

$$\partial_r = \partial_{r'}, \quad \partial_v = \partial_{v'}, \quad \partial_\theta = \partial_{\theta'} - ia \sin \theta (\partial_{v'} + \partial_{r'}), \quad \partial_\phi = \partial_{\phi'}.$$

Wir setzen diese Transformation in die Ausdrücke für die Nulltetrad ein. Dies ergibt Linearkombinationen der Vektoren $\partial_{v'}$, $\partial_{r'}$, $\partial_{\theta'}$ und $\partial_{\phi'}$ mit Koeffizienten in denen die Koordinaten r' und v' stehen. Wir fordern nun, dass diese reell sein sollen und setzen in den dann resultierenden Ausdrücken anstelle der gestrichenen Koordinaten wieder die ungestrichenen ein. Dies ergibt schließlich die neue Nulltetrad

$$\begin{aligned} \mathbf{l} &= \partial_r, \\ \mathbf{n} &= -\partial_v - \frac{1}{2} \left(1 - \frac{2mr}{r^2 + a^2 \cos^2 \theta} \right) \partial_r, \\ \mathbf{m} &= \frac{1}{\sqrt{2}(r + ia \cos \theta)} \left(-ia \sin \theta (\partial_r + \partial_v) + \partial_\theta + \frac{i}{\sin \theta} \partial_\phi \right). \end{aligned}$$

Dies ist wieder eine Basis mit der Eigenschaft, dass $\mathbf{l}^a$ und $\mathbf{n}^a$ reell sind und dass $\mathbf{m}^a$ und $\bar{\mathbf{m}}^a$ komplex konjugiert zueinander sind. Dies ist genau die Eigenschaft, die eine Nulltetrad besitzen muss. Wenn wir nun fordern, dass diese vier Vektoren eine solche Nulltetrad bilden sollen, dann wird dadurch eine Metrik definiert. Diese ist genau diejenige, bezüglich der die Vektoren die Skalarprodukte (10.3.1) besitzen.

Wir können nun die zu dieser Nulltetrad duale Basis des Kotangentialraums berechnen. Wir erhalten die vier 1-Formen

$$\begin{aligned} l_a &= -dv - a \sin^2 \theta d\phi, \\ n_a &= -\frac{1}{2} \alpha dv + dr + a \sin^2 \theta \left(1 - \frac{1}{2} \alpha \right) d\phi, \\ m_a &= -\frac{(r - ia \cos \theta)}{\sqrt{2}} (d\theta + i \sin \theta d\phi). \end{aligned}$$

Dabei ist

$$\alpha = 1 - \frac{2mr}{r^2 + a^2 \cos^2 \theta}.$$

Mit dieser dualen Basis kann man nun schließlich gemäß (10.3.2) die Metrik berechnen. Um verschiedene Koordinatensysteme unterscheiden zu können, benutzen wir ab jetzt Großbuchstaben für die Kerr-Eddington Koordinaten

$$g = \left(1 - \frac{2mR}{\Sigma}\right) dV^2 - 2dVdR + \frac{2mR}{\Sigma} (2a \sin^2 \Theta) dVd\Phi + 2a \sin^2 \Theta dRd\Phi - \Sigma d\Theta^2 - \left((R^2 + a^2) \sin^2 \Theta + \frac{2mR}{\Sigma} a^2 \sin^4 \Theta \right) d\Phi^2, \quad (10.3.3)$$

mit der Definition von $\Sigma = R^2 + a^2 \cos^2 \Theta$. Diese Metrik ist eine Lösung der Einstein-Gleichungen. Dies zu verifizieren ist für sich genommen schon eine sehr langwierige Aufgabe. Es handelt sich dabei um die Kerr-Metrik in einem Koordinatensystem, welches den avancierten Eddington-Finkelstein Koordinaten ähnelt (man beachte das Fehlen von dR^2 und den Term $-2dVdR$). Man nennt dieses Koordinatensystem **Kerr-Eddington Koordinaten**.

Kerr-Eddington
Koordinaten

Es ist offensichtlich, dass sich für $a = 0$ die Schwarzschild-Metrik ergeben muss, weil dann die Koordinatentransformation, die oben durchgeführt wurde, die identische Transformation ist. Dies lässt sich auch explizit nachprüfen. Gewissermaßen bekommt man die Kerr-Metrik, indem man die Schwarzschild-Metrik komplexifiziert, und dann eine Transformation ins Komplexe vornimmt. Die Kerr-Metrik ist also mit der Schwarzschild-Metrik durch eine komplexe Transformation verbunden. Die tiefere Bedeutung dieser Verwandtschaft ist noch nicht verstanden. Es gibt zwar ganz neue Untersuchungen von E.T.Newman zu diesem Thema¹, jedoch ist der letzte Durchbruch noch nicht gelungen.

10.3.2 Die Kerr-Metrik in Boyer-Lindquist Koordinaten

Die Kerr-Metrik enthält zwei Parameter, die Gesamtmasse m und den Drehimpuls $J = ma$, wobei a den spezifischen Drehimpuls darstellt. Die Kerr-Metrik kann in verschiedenen Koordinatensystemen aufgeschrieben werden. Je nach Wahl der Koordinaten werden unterschiedliche Bereiche der gesamten Kerr-Raumzeit überdeckt und verschiedene Aspekte verdeutlicht. Die Kerr-Metrik wurde mit dem **Spin-Formalismus** von Newman und Penrose gefunden, einem Formalismus, den wir hier nicht weiter erläutern wollen. Dies hatte jedoch zur Folge, dass sie zuerst in einem Koordinatensystem dargestellt wurde, in dem eine Interpretation dieser Geometrie sehr schwierig, wenn nicht sogar unmöglich war. Erst nachdem die **Boyer-Lindquist-Koordinaten** ge-

Spin-Formalismus

Boyer-Lindquist-
Koordinaten

¹Kozameh, Newman und Silva-Ortigoza, gr-qc/0506046

funden wurden, konnte man die Ähnlichkeit mit der Schwarzschild-Metrik erkennen und auch die physikalischen Eigenschaften der Metrik genauer untersuchen.

In diesen Koordinaten lautet die Metrik

$$g = dt^2 - \frac{2mr}{\Sigma} \left(dt - a \sin^2 \theta d\phi \right)^2 - \Sigma \left(\frac{1}{\Delta} dr^2 + d\theta^2 \right) - (r^2 + a^2) \sin^2 \theta d\phi^2. \quad (10.3.4)$$

Dabei sind die Funktionen Δ und Σ durch

$$\Delta = r^2 - 2mr + a^2, \quad (10.3.5)$$

$$\Sigma = r^2 + a^2 \cos^2 \theta \quad (10.3.6)$$

gegeben. Eine etwas andere Form dieser Metrik bekommt man, wenn man die Terme umstellt

$$g = \frac{\Delta}{\Sigma} \left(dt - a \sin^2 \theta d\phi \right)^2 - \Sigma \left(\frac{1}{\Delta} dr^2 + d\theta^2 \right) - \frac{1}{\Sigma} \sin^2 \theta \left((r^2 + a^2) d\phi - a dt \right)^2. \quad (10.3.7)$$

Diese Metrik sieht sehr kompliziert aus. Es ist jedoch in beiden Fällen sehr leicht zu sehen, dass man für $a = 0$ genau die Schwarzschild-Metrik zurück bekommt.

Die Transformation, die zwischen Kerr-Eddington-Koordinaten (V, R, Θ, Φ) und Boyer-Lindquist-Koordinaten (t, r, θ, ϕ) vermittelt, ist definiert durch

$$dV = dt + \frac{\Delta + 2mr}{\Delta} dr, \quad d\Phi = d\phi + \frac{a}{\Delta} dr, \quad dR = dr, \quad d\Theta = d\theta.$$

Übung 10.5: Bestimmen Sie die Beziehung zwischen den Koordinaten, die aus dieser Relation zwischen den Differentialen folgt.

Wir sehen, dass die Metrik offensichtlich für $\Sigma = 0$ singulär wird. Diese Singularität entspricht für $a \rightarrow 0$ der Schwarzschild-Singularität bei $r = 0$. Es scheint etwas eigenartig zu sein, dass

$$\Sigma = 0 \iff r = 0 \quad \text{und} \quad \theta = \frac{\pi}{2}$$

und nicht für $r = 0$ und weitere Werte von θ . Dies hat damit zu tun, dass die Koordinate r nicht genau das ist, was sie vorgibt zu sein. Insbesondere ist $r = 0$ nicht ein einzelner Punkt, sondern eine ganze Fläche. Wir werden später noch sehen, wie dies zustande kommt.

Eine weitere Singularität der Kerr-Metrik ist durch $\Delta = 0$ gegeben. Offensichtlich muß es sich dabei um ein Analogon zum Schwarzschild-Horizont bei $r = 2m$ handeln, denn für $a = 0$ ist $\Sigma/\Delta = 0 \iff r = 2m$. Jedoch ist es auch hier anders als bei Schwarzschild. Die Funktion $\Delta = r^2 - 2mr + a^2$ besitzt zwei Nullstellen

$$r = r_{\pm} = m \pm \sqrt{m^2 - a^2},$$

jedoch nur wenn $a^2 < m^2$ ist. Es handelt sich bei diesen Singularitäten der Metrik (10.3.7) offensichtlich um Koordinatensingularitäten, denn die Kerr-Metrik in Kerr-Eddington Koordinaten ist bei $R = r = r_{\pm}$ vollkommen regulär. Wir werden im Folgenden nur den Fall $a^2 < m^2$ betrachten. Der Fall $a = \pm m$ (extreme Kerr-Lösung) sowie der Fall $a^2 > m^2$ sind Sonderfälle, die eine weniger gute physikalische Interpretation besitzen.

10.3.3 Kerr-Schild Koordinaten

Das Koordinaten-System, in dem Kerr die Lösung entdeckte, fällt in die Klasse von Kerr-Schild Koordinaten, die sich gut eignen um sogenannte algebraisch spezielle Lösungen der Einstein-Gleichungen zu beschreiben. Wir erhalten diese Koordinaten, indem wir die Koordinatentransformation

$$\begin{aligned}x + iy &= (R - ia)e^{i\Phi} \sin \Theta, \\z &= R \cos \Theta, \\ \tau &= V - R.\end{aligned}\tag{10.3.8}$$

Mit dieser Transformation ergibt sich die Metrik in Kerr-Schild-Form

$$g = d\tau^2 - dx^2 - dy^2 - dz^2 - \frac{2mR^3}{R^4 + a^2z^2} \left(\frac{R(xdx + ydy) + a(xdy - ydx)}{R^2 + a^2} + \frac{zdz}{R} + d\tau \right)^2.\tag{10.3.9}$$

Dabei ist R wieder als Funktion von x, y und z zu betrachten, die durch die Gleichung

$$R^4 - (x^2 + y^2 + z^2 - a^2)R^2 - a^2z^2 = 0$$

definiert ist. Das wesentliche an dieser Form der Metrik ist die Tatsache, dass die flache Minkowski-Metrik als eine Art Hintergrund wirkt. Man kann dies benutzen, um das Verhalten der anderen Koordinatensysteme zu untersuchen, indem man beispielsweise die Flächen $R = \text{const}$, $\Theta = \text{const}$ usw. als Flächen in dem (zwar künstlichen) euklidischen Raum untersucht, der durch die 'kartesischen Koordinaten' (x, y, z) definiert wird.

Betrachten wir beispielsweise die Fläche $R = 0$, dann bekommen wir mit (10.3.8)

$$x = a \sin \Theta \sin \Phi, \quad y = a \sin \Theta \cos \Phi, \quad z = 0,$$

wobei $0 \leq \Phi \leq 2\pi$ und $0 \leq \Theta \leq \pi$. Diese Relationen beschreiben eine Scheibe mit Radius a , $\Theta = 0$ entspricht dem Zentrum der Scheibe, $\Theta = \pi/2$ ihrem Rand. Die Scheibe wird durch diese Beziehung zweimal überlagert, zu jedem Punkt auf der Scheibe gibt es zwei Θ -Werte.

Die Flächen mit konstantem $R \neq 0$ sind konfokale Ellipsoide, die rotationssymmetrisch um die z -Achse sind. Dies findet man schnell heraus, indem man die definierende Gleichung für R umschreibt in

$$\frac{x^2 + y^2}{R^2 + a^2} + \frac{z^2}{R^2} = 1.$$

Die Brennpunkte dieser Ellipsoide liegen auf einem Kreis mit Radius a . Für große Werte von R werden die Flächen immer kugelförmiger, so dass im Grenzfall R mit dem üblichen euklidischen Abstand zusammenfällt.

Die Flächen konstanten Θ 's lassen sich ebenfalls als Kegelschnitte darstellen. Elimination von R aus der Koordinatentransformation (10.3.8) liefert

$$\frac{x^2 + y^2}{a^2 \sin^2 \Theta} - \frac{z^2}{a^2 \cos^2 \Theta} = 1.$$

Es handelt sich also um Hyperboloide, rotationssymmetrisch um die z -Achse, die asymptotisch zu dem Kegel sind, der durch die Gleichung

$$x^2 + y^2 = z^2 \tan^2 \Theta$$

beschrieben wird. Diese Flächen sind in Abb. 10.5 etwas veranschaulicht. Es ist offensichtlich, dass die Metrik (10.3.9) die Form

$$g_{ab} = \eta_{ab} - h l_a l_b$$

besitzt. Es zeigt sich, dass für den Kovektor l_a die Gleichung $\eta_{ab} l^a l^b = 0$ gilt. Damit gilt auch $g_{ab} l^a l^b = l_a l^b = 0$. Es handelt sich also um einen lichtartigen Vektor. Im Fall der Kerr-Metrik ist

$$h = \frac{2mR^3}{R^4 + a^2 z^2}$$

und

$$l_a = d\tau + \frac{Rx + ay}{R^2 + a^2} dx + \frac{Ry - ax}{R^2 + a^2} dy + \frac{z}{R} dz.$$

Der Grenzfall $a = 0$ zeigt, dass auch die Schwarzschild-Metrik in der Kerr-Schild-Form geschrieben werden kann, mit

$$h = \frac{2m}{R}, \quad d\tau + \frac{x}{R} dx + \frac{y}{R} dy + \frac{z}{R} dz$$

10.3.4 Grundlegende Eigenschaften der Kerr-Metrik

Kontinuierliche Symmetrien Sowohl in der Kerr-Eddington-Form als auch in der Boyer-Lindquist-Form scheint die Kerr-Metrik nicht von der Zeit und vom Azimutwinkel abzuhängen. Die Metrik-Komponenten sind nur Funktionen von $r = R$ und $\theta = \Theta$. Die Transformationen

$$T \mapsto T + \alpha, \quad \text{bzw.} \quad t \mapsto t + \alpha,$$

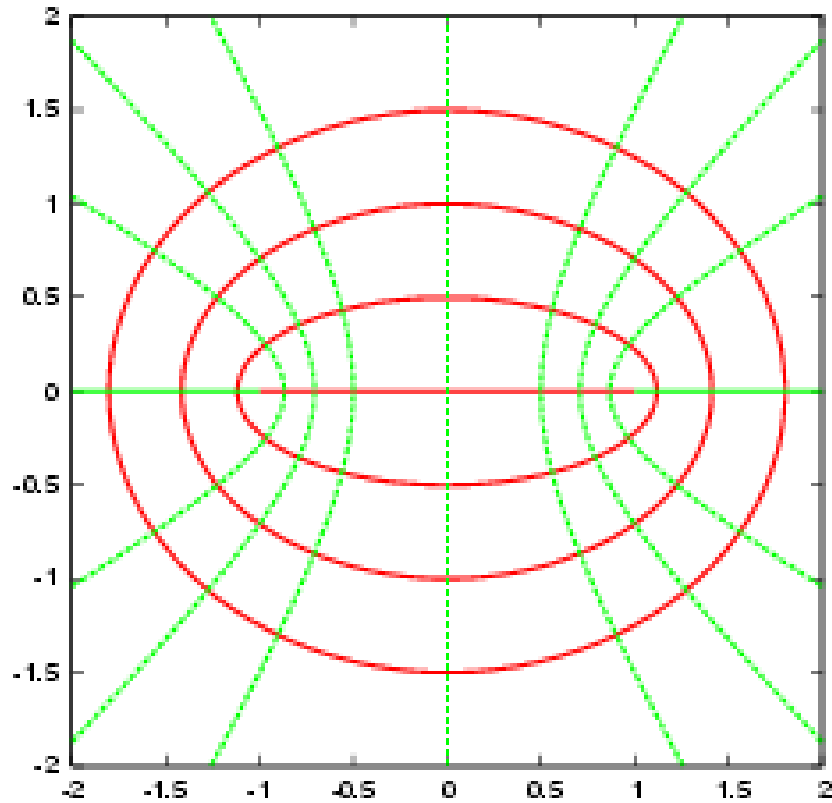


Abbildung 10.5: Die Schnitte der Flächen $R = \text{const}$ und $\Theta = \text{const}$ mit der y -Ebene.

die Zeittranslationen entsprechen, lassen die Metrik *invariant*, es sind dies also Symmetrie-Transformationen. Betrachten wir nur die Boyer-Lindquist-Form, dann ist $t \mapsto t + \alpha$ für beliebiges $\alpha \in \mathbb{R}$ eine *endliche* Symmetrie-Transformation. Betrachten wir einen beliebigen Punkt P auf der Kerr-Mannigfaltigkeit, der die Koordinaten (t, r, θ, ϕ) besitzt, dann ist durch $\alpha \mapsto t + \alpha$ eine *Kurve* durch P definiert. Deren Tangentialvektor T^α bei P ist durch die Ableitung nach α gegeben

$$T^\alpha[P] = \frac{d}{d\alpha}(t + \alpha, r, \theta, \phi) \Big|_{\alpha=0} = (1, 0, 0, 0)$$

gegeben. Auf diese Weise erhalten wir an jedem Punkt P einen Tangentialvektor $T^\alpha[P]$, mithin also ein Vektorfeld. Dieses Vektorfeld, welches man durch Differentiation einer Symmetrie-Transformation bekommt, nennt man auch *infinitesimale Symmetrie-Transformation* oder auch **Killing-Vektor**.

Killing-Vektor

Auf analoge Weise bekommen wir einen zweiten Killing-Vektor Φ^α , der infinitesimalen

Transformation in der ϕ -Richtung entspricht

$$\Phi^a[P] = \frac{d}{d\alpha}(t, r, \theta, \phi + \alpha) \big|_{\alpha=0} = (0, 0, 0, 1).$$

stationär
axial-symmetrisch

Die Kerr-Metrik ist demnach invariant unter Zeittranslationen und Rotationen um den Winkel ϕ . Sie ist **stationär** und **axial-symmetrisch**. In den Boyer-Lindquist Koordinaten haben die Killing-Vektoren die einfache Darstellung

$$T^a = \partial_t, \quad \Phi^a = \partial_\phi.$$

Wir können nun das Skalarprodukt dieser Vektoren berechnen

$$T_a T^a = g(\partial_t, \partial_t) = 1 - \frac{2m}{r} \left(\frac{r^2}{\Sigma} \right), \quad (10.3.10)$$

$$\Phi_a \Phi^a = g(\partial_\phi, \partial_\phi) = -r^2 \sin^2 \theta - a^2 \sin^2 \theta \left(1 + \frac{2mr \sin^2 \theta}{r^2 + a^2 \sin^2 \theta} \right) \quad (10.3.11)$$

Für $a = 0$ bekommen wir die Werte für die Schwarzschild-Metrik. Es zeigt sich, dass in diesem Fall der Betrag der infinitesimalen Zeittranslationen für große r positive ist, so wie man es auch im Minkowski-Raum hat. Translationen in der *Zeit* sollten entlang *zeitartigen* Richtungen geschehen. Jedoch wird der Killing-Vektor T^a in der Schwarzschild-Metrik wie auch in der Kerr-Metrik *raumartig* für kleine r .

Rotationen um eine Achse sollten raumartig sein. In der Tat ist $\Phi_a \Phi^a < 0$, sofern $r > 0$ bleibt. Aber muss das so sein? Wir hatten schon gesehen, dass $r = 0$ in der Kerr-Metrik eine kompliziertere Struktur hat, als wir es gewöhnt sind. Tatsächlich ist $r = 0$ (für eine feste Zeit t) nicht ein einzelner Punkt wie in der Minkowski-Raumzeit, sondern eine Scheibe, deren Rand der Kerr-Singularität entspricht. Nähert man sich der Scheibe von einer Seite aus in der Nähe ihres Zentrums, so wird r klein und erreicht $r = 0$. Nun kann man aber durch die Scheibe hindurch zu *negativen* Werten von r hindurch gehen. Diese Situation ist ähnlich wie bei Riemannschen Flächen, die durch geeignete Schnitte in verschiedene Blätter geteilt werden. Man kann über diese Schnitte in andere Blätter der Flächen gelangen. Hier kann man also ein Gebiet mit $r < 0$ erreichen, in dem dann der Killing-Vektor der infinitesimalen Rotationen *zeitartig* wird.

Diese Überlegungen zeigen, dass man in der Nähe der Kerr-Singularität mit Überraschungen rechnen muss. Beispielsweise ist die Transformation $\phi \mapsto \phi + \alpha$ periodisch, d.h., für $\alpha = 2\pi$ erhält man die Identitätsabbildung zurück. Die Kurven, die man für variables α bekommt sind geschlossen. In dem Bereich, in dem der Killing-Vektor Φ^a zeitartig ist, gibt es daher *geschlossene zeitartige Kurven*.

Diskrete Symmetrien Die Boyer-Lindquist-Form der Kerr-Lösung zeigt auch, dass die Metrik unter der simultanen diskreten Inversion von Zeit und Winkel

$$t \mapsto -t, \quad \phi \mapsto -\phi$$

invariant ist, jedoch nicht unter den separaten Transformationen. Dies deutet darauf hin, dass die Kerr-Lösung einen stationär rotierenden Körper beschreibt, denn in dieser Situation entspricht eine Umkehrung des Drehsinns genau einer Zeitinversion.

Ebenso ist die Metrik unter

$$t \mapsto -t, \quad a \mapsto -a, \quad \text{bzw.} \quad a \mapsto -a, \quad \phi \mapsto -\phi$$

invariant. Dies bedeutet, dass a den Drehsinn der Bewegung beschreibt, denn Umkehrung des Drehsinns muss durch Umkehrung von a kompensiert werden.

Schließlich ist das Auftreten des Diagonalterms mit $dt d\phi$ ebenfalls ein Hinweis auf die Interpretation als Feld eines rotierenden Körpers, denn dieser Term tritt auf, wenn man beispielsweise die flache Metrik in einem rotierenden Bezugssystem aufschreibt.

Übung 10.6: Zeigen Sie, dass die Minkowski-Metrik in einem System, welches sich bzgl. eines Inertialsystems um die z -Achse mit Winkelgeschwindigkeit a dreht, in Zylinderkoordinaten die Form

$$ds^2 = (1 - a^2 r^2) dt^2 - 2ar^2 d\phi dt - (dr^2 + r^2 d\phi^2 + dz^2)$$

annimmt.

Die statische Grenzfläche Die Boyer-Lindquist-Form der Kerr-Lösung geht für große r in die Minkowski-Metrik über. Das bedeutet, dass wir in diesem Grenzfall die Koordinaten (t, r, θ, ϕ) als Koordinaten in einem Inertialsystem im Unendlichen betrachten dürfen. Sie spezifizieren einen *Beobachter im Unendlichen*.

Wir betrachten nun einen zweiten Beobachter, der bzgl. des Inertialsystems im Unendlichen in Ruhe ist, das heißt, seine Ortskoordinaten (r, θ, ϕ) ändern sich nicht. Die Vierergeschwindigkeit u^a eines solchen *statischen Beobachters* ist proportional zum Koordinatenvektor ∂_t , es gilt also

$$u^a \sim \partial_t.$$

Eine Vierergeschwindigkeit ist immer ein *zeitartiger* Vektor, es muss daher immer

$$u_a u^a > 0$$

gelten. Nun ist aber

$$g_{tt} = g(\partial_t, \partial_t) = \Delta - a^2 \sin^2 \theta = r^2 - 2mr + a^2 \cos^2 \theta$$

und dieser Ausdruck verschwindet für

$$r = S_{\pm} = m \pm \sqrt{m^2 - a^2 \cos^2 \theta}. \quad (10.3.12)$$

Wie ist dies zu interpretieren? Wenn g_{tt} verschwindet, heißt das, dass ∂_t lichtartig wird. Da u^a proportional zu ∂_t sein soll, folgt daraus, dass der Beobachter zwangsläufig Lichtgeschwindigkeit erreicht. Man kann zeigen, dass die Weltlinie des Beobachters

keine Geodäte ist, das heißt ein solcher statischer Beobachter ist nicht frei fallend, sondern er muss dem Einfluß der Gravitation standhalten, indem er sich beschleunigt. Die Situation ist wie in einem Strudel, wo man gegen die Strömung schwimmen muss, um bzgl. dem Ufer in Ruhe zu sein. Je näher man an den Wirbel kommt, umso schneller muss man gegen die Strömung anschwimmen, d.h. die relative Geschwindigkeit zwischen Schwimmer und Wasser erhöht sich. Dasselbe passiert in der Nähe des Zentrums in der Kerr-Raumzeit: je näher der Beobachter dem Zentrum ist, umso stärker muss er beschleunigen, umso höher wird seine Geschwindigkeit, bis sie schließlich Lichtgeschwindigkeit erreicht. Ist dies der Fall, kann der Beobachter nicht mehr gegenüber dem Unendlichen in Ruhe bleiben. Aus diesem Grund nennt man die Fläche $r = S_+$ auch die **statische Grenzfläche**.

statische
Grenzfläche

Diese Fläche ist auch dadurch charakterisiert, dass Licht, welches von dieser Fläche ausgesandt wird, im Unendlichen eine *unendliche Rotverschiebung* erfährt. Dies ist ganz in Analogie zur Schwarzschild-Metrik, wo dies am Horizont $r = 2m$ geschieht. Offensichtlich ist $S_+ \rightarrow 2m$, $S_- \rightarrow 0$ für $a \rightarrow 0$. Im Unterschied zur Schwarzschild-Metrik ist die statische Grenzfläche jedoch kein Horizont, also keine lichtartige Fläche, die als eine Einweg-Membran wirken würde. Es ist also durchaus möglich, dass Beobachter, die diese Fläche ins Innere überschreiten, wieder herauskommen können. Die Membraneigenschaft des Schwarzschild-Horizonts wird von der Fläche $\Delta = 0$ übernommen, die dadurch ausgezeichnet ist, dass die Boyer-Lindquist-Form der Metrik dort singulär wird. Dieser sogenannte **Kerr-Horizont** ist durch die Gleichung

Kerr-Horizont

$$r = r_{\pm} = m \pm \sqrt{m^2 - a^2}$$

definiert.

Übung 10.7: Zeigen Sie, dass die statische Grenzfläche $r = S_+$ (außer für $\sin \theta = 0$) eine zeitartige Hyperfläche in der Kerr-Geometrie ist, dass aber der Horizont $r = r_+$ hingegen eine lichtartige Hyperfläche ist.

Auch hier gilt wieder $r_+ \rightarrow 2m$, $r_- \rightarrow 0$ für $a \rightarrow 0$, so dass im Schwarzschild-Limes der Schwarzschild-Horizont die charakterisierenden Eigenschaften der beiden Hyperflächen miteinander vereint.

Für $a > 0$ haben die statische Grenzfläche und der Horizont nur zwei Punkte gemeinsam, die Pole bei $\sin \theta = 0$. Ansonsten liegt die statische Grenzfläche immer außerhalb des Horizonts. Die Region, die von den beiden eingeschlossen wird, nennt man die **Ergosphäre**. Die Existenz dieser Region ermöglicht interessante Phänomene, wie wir im folgenden sehen werden. Die beiden Horizonte sowie die innere und äussere statische Grenzfläche sind in Abb. 10.6 veranschaulicht.

Ergosphäre

10.3.5 Der Penrose-Prozess

Die charakteristische Eigenschaft der Ergosphäre ist die Tatsache, dass in ihrem Inneren die Metrik-Funktion $g_{tt} = T_a T^a$ negativ wird, der Killing-Vektor $T^a = \partial_t$ wird raumar-



Abbildung 10.6: Horizonte und statische Grenzfläche im fast extremen Fall $a \approx m$

tig. Was im Unendlichen ein zeitartiger Vektor war, wird innerhalb der Ergosphäre ein raumartiger Vektor.

Um den Penrose-Prozess zu beschreiben, brauchen wir einen kurzen Exkurs über das Verhalten von frei fallenden Teilchen, d.h. von zeitartigen Geodäten. Diese werden durch eine Vierergeschwindigkeit u^a beschrieben, die der Geodäten-Gleichung

$$u^a \nabla_a u^b = 0$$

genügt. Ist ξ^a ein Killing-Vektor, also eine infinitesimale Symmetrie der Raumzeit, dann kann man zeigen, dass das Skalarprodukt

$$u_a \xi^a$$

ein *Konstante der Bewegung* ist, sich also entlang der Weltlinie des Teilchens nicht ändert.

Übung 10.8: Ein Killing-Vektor erfüllt die **Killing-Gleichung**

Killing-Gleichung

$$\nabla_{(a} \xi_{b)} = 0.$$

Zeigen Sie, dass aus dieser Gleichung und der Geodäten-Gleichung die obige Behauptung folgt.

In der Kerr-Geometrie haben wir zwei Killing-Vektoren, infinitesimale Zeittranslationen T^a und infinitesimale Drehungen Φ^a um die z -Achse. Es gibt daher zwei Konstanten der Bewegung eines Teilchens mit Ruhmasse m_0

$$E = m_0 u_a T^a, \quad L = u_a \Phi^a.$$

Dabei beschreibt E , welches mit der Zeittranslation verknüpft ist, die Energie des Teilchens, während L den Drehimpuls um die z -Achse angibt.

Wir betrachten nun ein Teilchen, welches im Unendlichen mit einer Energie E_0 startet und in die Ergosphäre eindringt. Wir nehmen weiter an, dass es dort in zwei Teile mit

Vierergeschwindigkeit v^a und Ruhmasse m_1 bzw. w^a und m_2 zerfällt. Aufgrund der Energieerhaltung muss

$$E_0 = m_1 v_a T^a + m_2 w_a T^a$$

gelten. Nun ist innerhalb der Ergosphäre T^a raumartig. Während für zwei zukunftsgerichtete zeitartige Vektoren T^a und u^a das Skalarprodukt $u_a T^a$ immer positiv ist, muss dies für raumartige Vektoren nicht gelten. Es gibt also die Möglichkeit, dass eines der Zerfallsprodukte (sagen wir Teilchen 1) mit negativer Energie ausgestattet wird und weiter ins Innere des Horizonts vordringt, während das zweite Teilchen mit positiver Energie die Ergosphäre wieder verläßt und ins Unendliche entschwindet.

Der Nettoeffekt dieses Prozesses ist, dass ein Teilchen mit Energie $E_2 > E_0$ im Unendlichen ankommt. Dem Kerr-Loch wurde Energie entzogen. Verfolgt man dieses Argument weiter und betrachtet auch den Drehimpuls, dann stellt man fest, dass diese Energie von der Rotationsenergie des Lochs herkommt. Man kann also im Prinzip ein rotierendes Kerr-Loch 'abbremsen' und ihm die Rotationsenergie entziehen.

Die Untersuchung der Schwarzen Löcher hat in der Vergangenheit auf viele faszinierende Ideen geführt: die Thermodynamik von schwarzen Löchern, der Flächensatz über das Anwachsen der Oberfläche von Schwarzen Löchern, Hawking Strahlung und das Verdampfen von schwarzen Löchern und schließlich in die Bereiche der Quanten-Gravitation. Die Untersuchung von schwarzen Löchern gibt uns Einblicke in die letzten, unverstandenen Gebiete der Physik.

A Differentialgeometrische Grundlagen

A.1 Mannigfaltigkeiten und Tangentialraum

Eine (differenzierbare) Mannigfaltigkeit M ist eine Menge von Punkten (im Zusammenhang mit ART auch Ereignisse genannt), die im Kleinen so aussieht wie ein $\mathbb{R}^n$. Die verschiedenen kleinen Teile werden durch Karten beschrieben, d.h., sie werden mit Koordinaten versehen. Karten, die die gleichen Bereiche überdecken, kann man aufeinander abbilden. Der Charakter der Übergangsabbildungen bestimmt den Charakter der Mannigfaltigkeit.

Man kann sich die Mannigfaltigkeit als ein Objekt vorstellen, welches durch 'Verkleben' von einzelnen Teilstücken konstruiert wird. Die 'Verklebevorschrift' wird durch Übergangsabbildungen gegeben. Abb. A.1 veranschaulicht diese Idee. Wir werden im

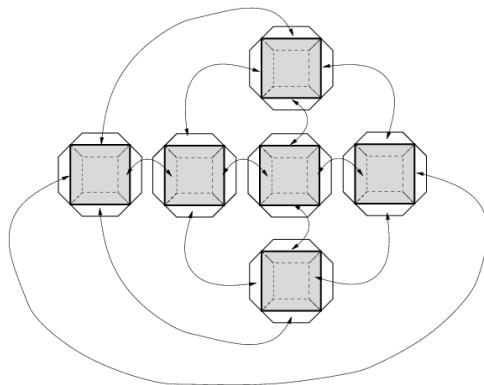


Abbildung A.1: Klebekonstruktion einer C^0 -Mannigfaltigkeit

folgenden nur differenzierbare Mannigfaltigkeiten betrachten, damit wir auf diesen Gebilden Analysis treiben können. D.h. wir wollen Ableitungen von (Skalar-, Vektor-, Tensor-) Feldern berechnen können.¹

¹Die Theorie der differenzierbaren Mannigfaltigkeiten ist eine zentrale Disziplin der Mathematik, die Differentialtopologie bzw. -geometrie. Daher gibt es unzählige Bücher. Es sei hier nur auf *Differential Topology* von V. Guillemin, A. Pollack verwiesen.

A.1.1 Mannigfaltigkeiten

Definition A.1. Eine n -dimensionale reelle Mannigfaltigkeit M ist ein topologischer (Hausdorff) Raum zusammen mit einer Familie von offenen Teilmengen U_α , die die folgenden Eigenschaften besitzt:

- (i) Die Mengen U_α überdecken M , d.h., jeder Punkt $P \in M$ liegt in mindestens einem U_α .
- (ii) Zu jedem α gibt es einen Homöomorphismus ψ_α , der U_α homöomorph auf eine offene Teilmenge des $\mathbb{R}^n$ abbildet.

Wir können uns also eine Mannigfaltigkeit zusammengesetzt denken aus kleinen Teilen U_α , die alle 'so aussehen wie ein kleines Stück des $\mathbb{R}^n$ '. Dies können beliebig viele sein.² Die offenen Mengen U_α heißen **Koordinatenumgebungen**, die Paare (U_α, ψ_α) nennt man **Karten**. Physiker verwenden anstelle der Bezeichnung 'Karte' üblicherweise '**Koordinatensystem**'. Ein Punkt $P \in M$ liegt in mindestens einer solchen Koordinatenumgebung U_α . Das Bild von P unter dem Homöomorphismus ψ_α , der U_α in den $\mathbb{R}^n$ abbildet, ist ein n -tupel von reellen Zahlen, den **Koordinaten** von P bzgl. der Koordinatenumgebung U_α , die wir mit $x^i[P]$ bezeichnen wollen, wobei $1 \leq i \leq n$.

Falls zwei Mengen U_α und U_β sich überlappen, also wenn $U_\alpha \cap U_\beta \neq \emptyset$ ist, dann ist die **Übergangsabbildung**

$$\psi_{\beta\alpha} = \psi_\beta \circ \psi_\alpha^{-1} : \psi_\alpha[U_\alpha \cap U_\beta] \rightarrow \psi_\beta[U_\alpha \cap U_\beta], \psi_\alpha(P) \mapsto \psi_\beta(P)$$

wohldefiniert und bijektiv. Jeder Punkt $P \in U_\alpha \cap U_\beta$ in der Übergangsregion zwischen zwei Koordinatenumgebungen U_α und U_β besitzt zwei Sätze von Koordinaten: $x^i[P]$ bzgl. U_α und $y^k[P]$ bzgl. U_β , die man eindeutig aufeinander abbilden kann (wegen der Bijektivität der Koordinatenabbildungen). Die Übergangsabbildung ist also von der Form $x^i = x^i(y^k)$, so dass

$$x^i[P] = x^i(y^k[P])$$

für alle $P \in U_\alpha \cap U_\beta$ gilt (vgl. Abb. A.2). Die Übergangsabbildung, auch **Koordinatentransformation** genannt, besteht aus n reellwertigen Funktionen von n reellen Variablen, die nach den üblichen Regeln der Differentialrechnung differenziert werden können. Die 'Qualität' der Übergangsabbildungen bestimmt die 'Qualität' der Mannigfaltigkeit.

Definition A.2. Die Mannigfaltigkeit heißt $\mathbb{C}^r$ -**differenzierbar (glatt, analytisch)**, falls die Übergangsabbildungen $x^i = x^i(y^k)$ r -mal stetig differenzierbar (unendlich oft differenzierbar, reell-analytisch) sind.

²Es gibt jedoch eine technische Zusatzvoraussetzung (oben unterdrückt), die verlangt, dass jeder Punkt nur in einer endlichen Anzahl von solchen Mengen liegen kann (parakompakt).

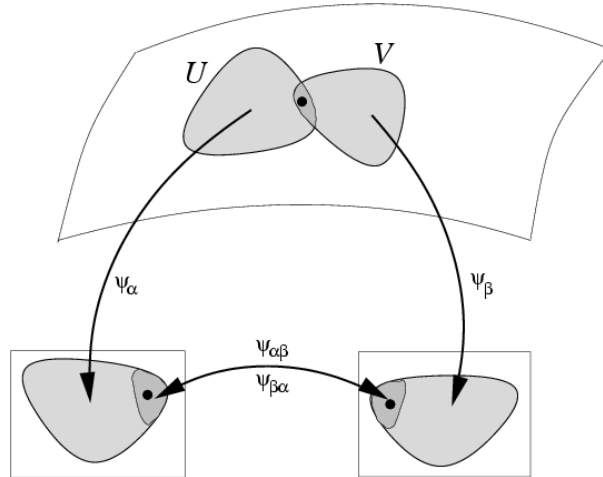


Abbildung A.2: Zur Definition einer Mannigfaltigkeit

Es ist üblich, in der Definition einer Mannigfaltigkeit anzunehmen, dass die Kollektion von offenen Mengen U_α maximal ist in dem Sinne, dass *alle* Koordinatensysteme mit kompatiblen Übergangsabbildungen eingeschlossen sind. Anderenfalls wäre es möglich, durch Hinzufügen oder Weglassen von Koordinatenumgebungen Mannigfaltigkeiten formal zu ändern, ohne dass dies wirklich sinnvoll wäre. Wir werden uns im Weiteren immer mit glatten Objekten (Mannigfaltigkeiten und Feldern) befassen.

Beispiele

1. $\mathbb{R}^n$ ist trivialerweise eine glatte Mannigfaltigkeit. Als (einzige) Karte können wir die Identitätsabbildung

$$U = \mathbb{R}^n \rightarrow \mathbb{R}^n, x \mapsto x$$

nehmen. (Man beachte, dass $\mathbb{R}^n$ eine offene Menge ist). Dadurch wird $\mathbb{R}^n$ zu einer n -dimensionalen glatten Mannigfaltigkeit.

2. Die 2-Sphäre $S^2 = \{(x, y, z) \in \mathbb{R}^3 : x^2 + y^2 + z^2 = 1\}$ ist, obwohl 2-dimensional, eine wichtige Mannigfaltigkeit, die an mehreren verschiedenen Stellen in der Physik auftritt. In der ART unter anderem als die Menge aller lichtartige Richtungen an einem Ereignis.

Wir setzen $U_1 = S^2 - \{(0, 0, 1)\}$ mit Koordinatenabbildung (vgl. Abb. A.3)

$$\psi_1 : U_1 \rightarrow \mathbb{R}^2; (x, y, z) \mapsto \xi = \left(\frac{x}{1-z}, \frac{y}{1-z} \right)$$

und $U_2 = S^2 - \{(0, 0, -1)\}$ mit Koordinatenabbildung

$$\psi_2 : U_2 \rightarrow \mathbb{R}^2; (x, y, z) \mapsto \eta = \left(\frac{x}{1+z}, \frac{y}{1+z} \right).$$

Dann rechnet man leicht nach, dass die Übergangsabbildung auf der Übergangsregion $U_{12} = S^2 - \{(0, 0, \pm 1)\}$ folgendermaßen lautet:

$$\eta = \eta(\xi) = \frac{\xi}{|\xi|^2}, \quad \text{wobei } |\xi|^2 = (\xi^1)^2 + (\xi^2)^2.$$

Diese Abbildung ist auf ihrem Definitionsbereich beliebig oft differenzierbar, sogar analytisch. Daher ist S^2 eine glatte und sogar analytische Mannigfaltigkeit.

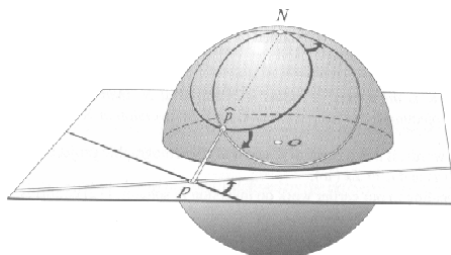


Abbildung A.3: Stereografische Projektion vom Nordpol

3. Der Kegel $C = \{(t, x, y, z) \in \mathbb{R}^4 : t^2 - x^2 - y^2 - z^2 = 0\}$ ist *keine* Mannigfaltigkeit.

Übung 1.1: Warum? Wie sieht es mit dem Halbkegel $C_+ = \{(t, x, y, z) \in C : t \geq 0\}$ aus? Und wie mit dem offenen Halbkegel $C_+^0 = \{(t, x, y, z) \in C : t > 0\}$?

A.1.2 Abbildungen zwischen Mannigfaltigkeiten

Definition A.3. Eine (stetige) Abbildung $f : M \rightarrow N$ zwischen zwei Mannigfaltigkeiten M und N heißt *differenzierbar in einem Punkt* $P \in M$, falls es Koordinatenumgebungen (U_α, ψ_α) bzw. (U_β, ψ_β) von P in M bzw. $f(P)$ in N gibt, so dass die Abbildung

$$f_{\alpha\beta} := \psi_\beta \circ f \circ \psi_\alpha^{-1}$$

differenzierbar in $x^i[P] = \psi_\alpha(P)$ ist.

Die Abbildung f ist *differenzierbar auf M* , falls sie in jedem Punkt differenzierbar ist.

Natürlich ist die Definition der Differenzierbarkeit in einem Punkt nicht abhängig von der Wahl der Karten. Wenn die Abbildung bzgl. eines Kartenpaares differenzierbar ist, dann auch bezüglich jedes anderen Paares.

Übung 1.2: Zeigen Sie, dass dies aus der Differenzierbarkeit der Übergangsabbildungen folgt.

Diffeomorphismus

Definition A.4. Eine differenzierbare Abbildung $f : M \rightarrow N$ heißt **Diffeomorphismus**, falls f bijektiv und auch f^{-1} differenzierbar ist.

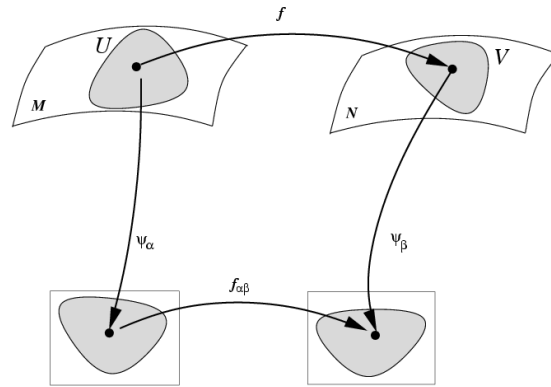


Abbildung A.4: Abbildung zwischen zwei Mannigfaltigkeiten

Beispiele

- (i) Rotationen der Kugel

$$R_\phi : S^2 \rightarrow S^2, \quad (x, y, z) \mapsto (x \cos \phi + y \sin \phi, -x \sin \phi + y \cos \phi, z).$$

- (ii) Die Inversion oder Antipodenabbildung der Sphäre

$$i : S^2 \rightarrow S^2, (x, y, z) \mapsto (-x, -y, -z).$$

- (iii) Ein bekanntes Beispiel für eine bijektive, differenzierbare Abbildung, die kein Diffeomorphismus ist, ist

$$f : \mathbb{R} \rightarrow \mathbb{R}, x \mapsto x^3.$$

Übung 1.3: Verifizieren Sie diese Beispiele.

Von besonderem Interesse (besonders in der ART) sind die Diffeomorphismen einer Mannigfaltigkeit auf sich. Sie bilden eine Gruppe, die **Diffeomorphismengruppe** bzgl. der Komposition $\circ$ von Abbildungen.

Diffeomorphismengruppe

A.2 Felder

A.2.1 Skalarfelder

Ein (reelles) **Skalarfeld** ist eine differenzierbare Abbildung $\phi : M \rightarrow \mathbb{R}$. Bezüglich einer Karte (U_α, ψ_α) hat es die Darstellung $\phi_\alpha = \phi \circ \psi_\alpha^{-1}$, die durch

$$\phi_\alpha(x^i[P]) = \phi(P)$$

für alle $P \in U_\alpha$ definiert ist.

Folglich entspricht jedem Skalarfeld eine Kollektion (ϕ_α) von reellwertigen Funktionen von n Variablen, die miteinander verträglich sind in dem Sinne, dass für alle $P \in U_\alpha \cap U_\beta$ gilt

$$\phi_\alpha(x^i[P]) = \phi(P) = \phi_\beta(y^k[P]) \Rightarrow \phi_\alpha(x^i) = \phi_\beta(y^k(x^i)).$$

Etwas anders geschrieben, lautet die Verträglichkeitsbedingung

$$\phi_\alpha = \phi_\beta \circ (\psi_\beta \circ \psi_\alpha^{-1}).$$

Stellen wir uns umgekehrt vor, wir haben eine Kollektion von Funktionen (ϕ_α) gegeben, die auf den entsprechenden Koordinatenumgebungen definiert sind. Wir fragen uns, wann diese 'Teilfelder' zu einem eindeutigen Skalarfeld auf ganz M 'zusammengeklebt' werden können, welches dann also 'global' definiert ist. Dies ist offensichtlich genau dann der Fall, wenn die obige Kompatibilitätsbedingung gilt, denn nur dann ist garantiert, dass das Skalarfeld auf den Überlappregionen wohldefiniert ist. Sind die lokalen Teilfelder alle differenzierbar, so ist auch das global zusammengeklebte Skalarfeld per definitionem differenzierbar.

Wir können also Skalarfelder als global definierte Funktionen betrachten, oder sie als eine Kollektion von lokal definierten, aber miteinander kompatiblen Koordinatendarstellungen auffassen. Dieser „patch-work“ Zugang zur Definition von globalen Objekten stellt sich manchmal als recht vorteilhaft heraus und wir werden ihn dazu verwenden, um Vektorfelder bzw. Tensorfelder im allgemeinen zu definieren.

Ist $f : M \rightarrow N$ eine differenzierbare Abbildung und ϕ ein Skalarfeld auf N , dann ist $\phi \circ f$ ein Skalarfeld auf M . Wir nennen dieses „verpflanzte“ Skalarfeld den **pull-back** von ϕ unter f und bezeichnen es mit $f^*\phi$.

A.2.2 Kontravariante Vektorfelder

Mit dem Begriff 'Vektor' verknüpft man intuitiv eine Richtung. Für einen Vektor im $\mathbb{R}^n$, also ein n -tupel $(V^1, \dots, V^n)$ kann man diese Intuition sogar präzisieren, indem man das n -tupel mit der *Richtungsableitung* einer beliebigen Funktion f in Richtung des Vektors in Verbindung bringt:

$$(V^1, \dots, V^n) \longleftrightarrow \sum_i V^i \frac{\partial f(x)}{\partial x^i}$$

Mit Hilfe jedes fest gewählten n -tupels kann man für jede Funktion f die Richtungsableitung berechnen. Umgekehrt sind auch alle V^i bestimmbar, wenn man für jede Funktion weiß, wie ihre Richtungsableitung aussieht. Dem Vektor $(V^1, \dots, V^n)$ entspricht also in eindeutiger Weise der lineare Differentialoperator $\sum_i V^i \frac{\partial}{\partial x^i}$.

Dies lässt sich auf Mannigfaltigkeiten übertragen. Man betrachtet zunächst lineare Differentialoperatoren in einer Karte. In einer Karte (U_α, ψ_α) hat ein linearer Differentialoperator am Punkt $x^i[P] \in \psi_\alpha[U_\alpha]$ die Darstellung

$$\sum_i V_\alpha^i(x) \frac{\partial}{\partial x^i},$$

wobei $\frac{\partial}{\partial x^i}$ die partielle Ableitung nach der Koordinate x^i bedeutet. Befindet sich P in einem Übergangsgebiet $U_{\alpha\beta} = U_\alpha \cap U_\beta$ zwischen zwei Karten U_α und U_β , dann gibt es auch in der Karte (U_β, ψ_β) lineare Differentialoperatoren. Wir haben also zwei mögliche Darstellungen

$$\sum_i V_\alpha^i(x) \frac{\partial}{\partial x^i} \quad \text{und} \quad \sum_k V_\beta^k(y) \frac{\partial}{\partial y^k}.$$

Wann stimmen diese überein? Wir betrachten ein Skalarfeld ϕ auf $U_{\alpha\beta}$, so dass $\phi_\alpha(x) = \phi_\beta(y(x))$. Dann sind die beiden Ausdrücke gleich, falls für jedes solche Skalarfeld die beiden Ableitungen für alle $P \in U_{\alpha\beta}$ überein stimmen:

$$\sum_i V_\alpha^i(x) \frac{\partial \phi_\alpha}{\partial x^i}(x) = \sum_k V_\beta^k(y) \frac{\partial \phi_\beta}{\partial y^k}(y). \quad (\text{A.2.1})$$

Nun gilt aber

$$\begin{aligned} \sum_i V_\alpha^i(x) \frac{\partial \phi_\alpha}{\partial x^i}(x) &= \sum_i V_\alpha^i(x) \frac{\partial \phi_\beta(y(x))}{\partial x^i} \\ &= \sum_{ik} V_\alpha^i(x) \frac{\partial \phi_\beta}{\partial y^k}(y(x)) \frac{\partial y^k}{\partial x^i}(x) \\ &= \sum_k V_\beta^k(y) \frac{\partial \phi_\beta}{\partial y^k}(y). \end{aligned}$$

Die beiden Ausdrücke sind also genau dann gleich, falls

$$\sum_i V_\alpha^i(x) \frac{\partial y^k}{\partial x^i}(x) = V_\beta^k(y(x)). \quad (\text{A.2.2})$$

Betrachten wir nun alle Koordinatenumgebungen U_α , in denen sich P befindet. Wenn wir in einer Karte ein n -tupel $(V^1(x^i[P]), \dots, V^n(x^i[P]))$ angeben, dann kann in jeder anderen Karte ein n -tupel berechnet werden, indem man von dem Transformationsgesetz Gebrauch macht. Wir haben so eine Kollektion von n -tupeln konstruiert, je eines für jede Karte, die miteinander kompatibel sind. Offensichtlich wird durch diese Kollektion ein kartenunabhängiges Objekt definiert, denn berechnen wir die Ableitung einer beliebigen Funktion in P bzgl. einer Karte, dann bekommen wir genau das gleiche Ergebnis, wie wenn wir es bzgl. einer anderen Karte berechnet hätten. Dieses 'Objekt' ist in jeder Karte eine Richtungsableitung. Wir kommen so zu der

Definition A.5. Sei $P \in M$ und U_α die Kollektion von Koordinatenumgebungen mit $P \in U_\alpha$. Eine Kollektion von n -tupeln $(V_\alpha^1, \dots, V_\alpha^n)$ definiert einen **Tangentialvektor** an M im Punkt P , falls die n -tupel miteinander kompatibel im Sinne von (A.2.2) sind.

Tangential

Diese Definition ist unanschaulich. Es wäre einfacher gewesen, eine Mannigfaltigkeit als eine Teilmenge in einem $\mathbb{R}^n$ zu veranschaulichen, den Vektorbegriff im $\mathbb{R}^n$ zu verwenden und so einen Tangentialvektor zu definieren. Zumindest in drei Dimensionen hätte man dann eine klare Vorstellung von den Verhältnissen gehabt. Diese Vorgehensweise ist aber zum einen unbefriedigend, weil man von etwas ($\mathbb{R}^n$) Gebrauch macht, was nicht streng notwendig ist. Andererseits ist es aber auch physikalisch falsch, anzunehmen dass es einen umgebenden Raum für die Raumzeit gäbe. Denn worin sollte sich unser Universum befinden? Sollten wir nicht in der Lage sein, nur mit den Mitteln, die wir zur Verfügung haben, einen Vektor zu definieren, also ohne auf die Struktur eines uns sowieso nicht zugänglichen umgebenden Raumes zurück zu greifen.

Dennoch ist die Definition im Grunde einfach. Ein Tangentialvektor an einem Punkt wird in einem Koordinatensystem als ein n -tupel dargestellt und durch den 'Kunstgriff' alle möglichen kompatiblen Darstellungen als gleichberechtigt in die Beschreibung aufzunehmen, wird die Definition koordinatenunabhängig. Es gibt andere Definitionen eines Tangentialvektors, die andere Eigenschaften eines Vektors im $\mathbb{R}^n$ auf Mannigfaltigkeiten übertragen (Tangentialvektor an eine Kurve, Produktregel der Richtungsableitung etc.) Sie alle haben die Eigenschaft, dass eine Klasse von konkreten 'Darstellungen' durch eine Äquivalenzrelation zu einem abstrakten Objekt zusammengefasst wird. Letzten Endes sind alle diese Definitionen äquivalent.

Offensichtlich bilden die Tangentialvektoren an jedem Punkt P einen linearen Vektorraum der Dimension n . Denn in einer Karte ist das so (die Menge aller n -tupel ist ein n -dimensionaler Vektorraum) und die Linearität der Kompatibilitätsbedingung (A.2.2) hat zur Folge, dass die Linearkombination zweier Tangentialvektoren kartenunabhängig definiert ist. Dieser Vektorraum heißt **Tangentialraum** an M in P und wird mit $T_P M$ bezeichnet.

Tangentialraum

Koordinatenbasis

In jedem Koordinatensystem gibt es eine ausgezeichnete Basis, die **Koordinatenbasis**. Sie ist dadurch definiert, dass wir in einer festen Karte (U_α, ψ_α) mit Koordinaten x^i die n -tupel $(1, 0, \dots, 0), (0, 1, 0, \dots, 0), \dots, (0, \dots, 0, 1)$ wählen. Dies definiert n Tangentialvektoren in P , die linear unabhängig sind. Sie werden üblicherweise mit

$$\frac{\partial}{\partial x^i} \Big|_P$$

bezeichnet, denn es handelt sich um Differentialoperatoren, die die Änderung eines Skalarfeldes ϕ am Punkt P entlang einer Koordinatenlinie liefern. Eine solche Basis ist mit jeder Karte verknüpft, d.h. jedes Koordinatensystem definiert eine Basis in $T_P M$. Diese Basen müssen sich natürlich ineinander transformieren lassen, d.h. sind $(\frac{\partial}{\partial x^i} \Big|_P)_{i=1:n}$ und

$(\frac{\partial}{\partial y^k} \big|_P)_{k=1:n}$ zwei Koordinatenbasen in $T_P M$, dann gibt es n^2 reelle Zahlen $s^i_k(P)$, so dass

$$\frac{\partial}{\partial y^k} \big|_P = \sum_i s^i_k(P) \frac{\partial}{\partial x^i} \big|_P$$

Übung 1.4: Wodurch sind diese Koeffizienten bestimmt?

Ein beliebiger Tangentialvektor $V \in T_P M$ lässt sich nun bzgl. dieser beiden Basen darstellen

$$\sum_i V^i(P) \frac{\partial}{\partial x^i} \big|_P = V = \sum_k \hat{V}^k(P) \frac{\partial}{\partial y^k} \big|_P.$$

Setzen wir hier die Beziehung zwischen den Basen ein, dann bekommen wir das bekannte Transformationsverhalten, vgl. (3.1.2), für die Koordinaten eines kontravarianten Vektors³.

$$V^i(P) = \sum_k s^i_k(P) \hat{V}^k(P), \quad i = 1 : n.$$

Wir haben nun die schon früher angesprochene Situation, dass sich an jedem Punkt der Mannigfaltigkeit ein Vektorraum, nämlich der Tangentialraum, befindet. All diese Räume haben im allgemeinen nichts miteinander zu tun. Jeder Tangentialvektor ist automatisch mit einem Fußpunkt P versehen und Tangentialvektoren mit verschiedenen Fußpunkten liegen in verschiedenen Vektorräumen, sie können also nicht addiert werden. Man kann sich diese Situation (etwas naiv, aber hilfreich) vorstellen, als ob an jedem Punkt der Mannigfaltigkeit eine Faser angeheftet wäre, die aus allen Vektoren aus $T_P M$ besteht. Die Menge all dieser Fasern besteht dann aus allen möglichen Tangentialvektoren an allen möglichen Punkten der Mannigfaltigkeit. Sie hat eine spezielle Struktur: man kann von jedem ihrer Elemente, also von jedem Tangentialvektor sagen, an welchem Punkt er angeheftet ist. Diese Vorstellung lässt sich formalisieren.

Die Vereinigung $TM := \bigcup_{P \in M} \{P\} \times T_P M$ heißt **Tangentialbündel** von M . Das Tangentialbündel ist auf natürliche Weise ebenfalls eine Mannigfaltigkeit. Um dies zu sehen müssen wir Karten angeben. Seien x^i lokale Koordinaten für eine Umgebung U_α von M . Dann lässt sich ein Tangentialvektor an einen Punkt P in U_α eindeutig als Linearkombinationen der Koordinatenableitungen darstellen:

Tangentialbündel

$$V = \sum_i V^i \frac{\partial}{\partial x^i}.$$

Wir können also eine Koordinatenumgebung $V_\alpha := U_\alpha \times \mathbb{R}^n \subset TM$ definieren, sowie einen Homöomorphismus $\phi_\alpha : V_\alpha \rightarrow \mathbb{R}^{2n}$, der jedem Tangentialvektor $V \in T_P M$ die Koordinaten (x^i, V^k) zuweist. Dies definiert eine Karte von TM , eine sogenannte **Bündelkarte**. Mit jeder Karte U_α von M ist eine Bündelkarte verknüpft, die Bündelkarten

Bündelkarte

³Dies darf uns nicht verwundern, denn dieses Verhalten ist allgemein gültig für jeden Vektor in jedem Vektorraum.

überdecken TM . Die Übergangsabbildungen dieser Karten sind kompatibel und verleihen damit TM eine Mannigfaltigkeitsstruktur.

Der Name „Bündel“ kommt daher, dass an jedem Punkt von M ein Tangentialraum, die *Faser*, angeheftet ist. Die Tatsache, dass jede Faser zu einem Punkt gehört wird durch die Projektionsabbildung $\pi : TM \rightarrow M$ ausgedrückt, die jedem Tangentialvektor aus TM seinen „Fußpunkt“ zuordnet. Diese Abbildung ist differenzierbar, denn in einer Bündelkarte V_α und assoziierter Karte U_α von M ist die Projektionsabbildung gegeben durch

$$(x^i, V^k) \mapsto x^i.$$

Vektorfeld Ein **Vektorfeld** ist eine differenzierbare Abbildung $V : M \rightarrow TM$ mit der Eigenschaft, dass $\pi \circ V = \text{id}_M$ ist. Das bedeutet, dass jedem Punkt P ein Element aus der Faser bei P , also ein Tangentialvektor aus $T_P M$, zugeordnet wird. Man sagt auch, ein Vektorfeld sei ein **Schnitt** im Tangentialbündel.

Bemerkung Es gibt andere Möglichkeiten, Tangentialvektoren, das Tangentialbündel und Vektorfelder zu definieren. Jede hat ihre Vor- und Nachteile. Das wesentliche am obigen Zugang ist die Tatsache, dass man lokale Definitionen ‘globalisiert’, indem man Kompatibilitätstransformationen fordert.

Differential
Tangentialabbildung

Eine Abbildung $f : M \rightarrow N$, die bei P differenzierbar ist, induziert eine Abbildung $f_*(P) : T_P M \rightarrow T_{f(P)} N$, genannt **Differential** oder **Tangentialabbildung** von f in P . Diese wird auch mit $T_P f$ bezeichnet und ist folgendermaßen definiert: sei $y^k = f^k(x^i)$ die Koordinatendarstellung von f bzgl. eines Kartenpaares um P in M bzw. $f(P)$ in N und V^i die Komponenten eines Tangentialvektors in $T_P M$, dann setzen wir

$$(f_*(P)(V))^k = \sum_i V^i \frac{\partial f^k}{\partial x^i}(x[P])$$

als die Komponenten eines Tangentialvektors bei $f(P)$.

Übung 1.5: Prüfen Sie nach, dass die Transformationseigenschaften unter Kartenwechsel in M erfüllt sind.

Eine andere, mehr geometrische Art, diese Abbildung zu definieren ist die folgende. Wir betrachten eine Kurve $\mathbb{R} \rightarrow M$, $t \mapsto \gamma(t)$ mit $\gamma(0) = P$. Diese definiert einen Tangentialvektor $V = \dot{\gamma}(0)$ in $T_P M$, indem wir einfach bzgl. einer Karte $V^i = \frac{d}{dt} x^i[\gamma(t)]|_{t=0}$ setzen. Nun betrachten wir die Bildkurve von γ unter f . Dies ist eine Kurve $(f \circ \gamma)(t)$ in N mit $(f \circ \gamma)(0) = f(P)$. Sie definiert einen Tangentialvektor $f_*(P)(V)$ in $f(P)$ auf analoge Weise wie γ . In lokalen Koordinaten wie oben bekommt man wieder

$$\begin{aligned} \frac{d}{dt} y^k[(f \circ \gamma)(t)]|_{t=0} &= \frac{d}{dt} f^k(x(\gamma(t)))|_{t=0} \\ &= \sum_i \frac{\partial f^k}{\partial x^i}(x[P]) \frac{d}{dt} x^i[\gamma(t)]|_{t=0} = \sum_i V^i \frac{\partial f^k}{\partial x^i}(x[P]). \end{aligned}$$

Die Tangentialabbildung der Abbildung f ist in Koordinaten demnach nichts anderes als die Anwendung der Funktional- (oder Jacobi-) Matrix auf die Komponenten eines Tangentialvektors.

Wenn f überall differenzierbar ist, dann gibt es auch an jedem Punkt die Differentialabbildung $f_*(P)$ und wir können insgesamt eine Abbildung $f_* : TM \rightarrow TN, V_P \mapsto f_*(P)V_P$ der entsprechenden Tangentialbündel definieren. Mithilfe von Bündelkarten zeigt man leicht, dass sie differenzierbar ist.

Mithilfe von f lassen sich Vektorfelder von M nach N transportieren: ist V ein Vektorfeld auf M , dann ist f_*V ein Vektorfeld auf N . Man nennt diese Operation auch **push-forward**.

push-forward

A.2.3 Kovariante Vektorfelder

Eine spezielle Sorte von Abbildungen zwischen Mannigfaltigkeiten sind Skalarfelder, deren Tangential-Abbildung wir nun anschauen. Es sei also $\phi : M \rightarrow \mathbb{R}$ eine reellwertige Funktion auf M . An jedem Punkt P ist das Differential von ϕ eine Abbildung $T_P\phi : T_P M \rightarrow T_{\phi(P)}\mathbb{R} = \mathbb{R}$ eine reellwertige Abbildung des Tangentialraums an M in P , also eine Linearform auf $T_P M$. Was ist die Bedeutung dieser Linearform? Sei $V \in T_P M$ ein beliebiger Tangentialvektor und es sei $\gamma : \mathbb{R} \rightarrow M$ eine Kurve mit $\gamma(0) = P$ und $\dot{\gamma}(0) = V$, dann ist

$$\langle T_P\phi, V \rangle = T_P\phi(V) = \frac{d\phi(\gamma(t))}{dt} \Big|_{t=0} = \sum_i V^i \frac{\partial \phi}{\partial x^i}(x[P]),$$

wobei die letzte Gleichung im Bereich einer Karte gilt. Diese Formel erlaubt uns eine anschauliche Interpretation: die Linearform $T_P\phi$ gibt die *infinitesimale Änderung von ϕ in Richtung der Kurve γ* an, also die Richtungsableitung von ϕ in Richtung V . Man bezeichnet $T_P\phi$ für Skalarfelder auch mit $d\phi_P$ und nennt die Abbildung $d\phi$ das (totale) Differential von ϕ .

Wir sehen also, dass es an jedem Punkt P einer Mannigfaltigkeit zum Tangentialraum $T_P M$ auch einen Dualraum $T_P^* M$ gibt⁴, den Raum der Linearformen auf $T_P M$, und dass die Differentiale von Skalarfeldern eben solche Linearformen liefern.

Jedes Koordinatensystem bei P liefert eine ausgezeichnete Basis von $T_P M$, die Koordinatenbasis. Zu dieser wiederum gibt es aufgrund der allgemeinen Theorie eine duale Basis im Dualraum $T_P^* M$. Kann man diese auch durch die Koordinaten beschreiben? Betrachten wir die Basisvektoren

$$\left. \frac{\partial}{\partial x^i} \right|_P$$

⁴Dies ist natürlich offensichtlich, denn jeder endlich-dimensionale Vektorraum besitzt einen Dualraum gleicher Dimension.

bzgl. eines Koordinatensystems $(x^i)_{i=1:n}$. Die duale Basis ω^i , $i = 1 : n$ in T_p^*M ist charakterisiert durch die Dualitätsrelation

$$\left\langle \omega^i, \frac{\partial}{\partial x^k} \Big|_P \right\rangle = \delta_k^i.$$

Wir konstruieren nun n Funktionen auf der Koordinatenumgebung U von P , deren Differentiale bei P genau diese Eigenschaft besitzen. Dazu betrachten wir die Abbildungen $\pi^i : U \rightarrow \mathbb{R}$, die jedem Punkt $Q \in U$ in der Koordinatenumgebung seine i -te Koordinate zuordnen. Dies sind genau n skalare Funktionen.

Die Kurve $\gamma : \mathbb{R} \rightarrow M$, die dadurch definiert ist, dass ihre Koordinatendarstellung durch

$$t \mapsto (x^1, \dots, x^k + t, \dots, x^n)$$

gegeben ist, hat als Tangentialvektor in P gerade den Basisvektor $\frac{\partial}{\partial x^k} \Big|_P$ und es gilt

$$\pi^i(\gamma(t)) = \begin{cases} x^i, & i \neq k \\ x^k + t, & i = k \end{cases}.$$

Damit erhalten wir

$$\left\langle d\pi^i, \frac{\partial}{\partial x^k} \Big|_P \right\rangle = \frac{d\pi^i(\gamma(t))}{dt} \Big|_{t=0} = \delta_k^i.$$

Die duale Basis einer Koordinatenbasis $\frac{\partial}{\partial x^i} \Big|_P$ ist demnach gegeben durch die **Koordinatendifferentiale** $d\pi^i$, die man allgemein etwas laxer als dx^i schreibt.

Koordinaten-
differentiale

Das totale Differential eines Skalarfeldes ϕ lässt sich daher im Bereich einer Karte als Linearkombination

$$d\phi = \sum_i \alpha_i dx^i$$

schreiben. Die Koeffizienten bestimmen wir unter Verwendung der Dualitätsrelation, indem wir beide Seiten auf einen Koordinatenbasisvektor anwenden.

Übung 1.6: Zeigen Sie, dass dann

$$d\phi = \sum_i \frac{\partial \phi}{\partial x^i} dx^i \tag{A.2.3}$$

gilt.

Bemerkung In der Physik wird (A.2.3) dahingehend interpretiert, dass sich die infinitesimale Änderung $d\phi$ von ϕ aus den infinitesimalen Änderung der Koordinaten mithilfe der partiellen Ableitungen berechnen lässt. Man beachte jedoch, dass hier nichts infinitesimal wird. Vielmehr sind die Differentiale (im Gegensatz zu Differenzen) nicht mehr auf der Mannigfaltigkeit definiert, sondern im Tangentialraum eines Punktes. Dies ist im übrigen eine allgemeine Regel. Wenn in der Physik Verhältnisse in der Nähe eines Punktes 'in erster Ordnung' bzw. 'infinitesimal' betrachtet werden, dann bedeutet

dies nichts anderes, als dass man in dem jeweiligen Punkt die Mannigfaltigkeit durch ihren Tangentialraum in dem Punkt ersetzt und entsprechende Abbildungen durch ihre Tangentialabbildung.

In gleicher Weise wie man die Tangentialräume an eine Mannigfaltigkeit zum Tangentialbündel zusammenfasst, kann man auch die entsprechenden Dualräume zu einem Bündel zusammenfassen, dem **Kotangentenbündel** $T^*M = \bigcup_{P \in M} \{P\} \times T_P^*M$.

Kotangentenbündel

Ebenso wie dem Tangentialbündel kann man dem Kotangentenbündel eine Mannigfaltigkeitsstruktur geben. Die entsprechenden Karten sind analog zu den Bündelkarten für TM konstruiert. Sind x^i lokale Koordinaten für eine Umgebung U von M , dann lässt sich jeder Kotangentenvektor ω an einem Punkt $P \in U$ in der Form

$$\omega = \sum_k \omega_k dx^k$$

als Linearkombination der Koordinatendifferentiale darstellen. Das $2n$ -tupel (x^i, ω_i) liefert also eine Koordinatendarstellung für alle Kotangentenvektoren über U . Wieder zeigt man leicht, dass die so definierten Karten kompatibel sind.

Analog zu Vektorfeldern definieren wir ein kovariantes Vektorfeld oder **Kovektorfeld** als Schnitt in T^*M , also eine Abbildung $\omega : M \rightarrow T^*M$, die jedem Punkt P in M ein Element aus dem Kotangentenraum bei P zuordnet. Man nennt ein kovariantes Vektorfeld auch **1-Form**(feld).

Kovektorfeld

1-Form

Übung 1.7: Bestimmen Sie das Verhalten der Koordinatendarstellung bei Kartenwechsel und rechtfertigen Sie so die Bezeichnung 'kovariant'.

Analog dem push-forward für Vektorfelder, gibt es auch für 1-Formen eine Möglichkeit, sie von Mannigfaltigkeit zu Mannigfaltigkeit zu transportieren, den **pull-back** oder auch **Zurückziehen**. Wie der Name andeutet, ist die Richtung dieser Abbildung entgegengesetzt zu $f : M \rightarrow N$. Ist ω eine 1-Form auf N , dann wird der pull-back $f^*\omega$ auf M durch seine Wirkung auf beliebige Vektorfelder X auf M definiert: 3

pull-back

Zurückziehen

$$\langle f^*\omega, X \rangle(P) = \langle \omega, f_*X \rangle(f(P))$$

Betrachten wir als speziellen Fall das Differential eines Skalarfeldes ϕ auf M . Sind y^k bzw. x^i lokale Koordinaten in N bzw. M , und $y^k = f^k(x^i)$ die Koordinatendarstellung von f , dann gilt für jedes Vektorfeld $V = V^i \partial / \partial x^i$ auf M

$$\begin{aligned} \langle d(f^*\phi), V \rangle(x) &= \langle d(\phi \circ f), V \rangle(x) = \sum_{ik} \left(\frac{\partial \phi}{\partial y^k}(y) \frac{\partial f^k}{\partial x^i}(x) \right) \langle dx^i, V \rangle(x) \\ &= \sum_{ik} \frac{\partial \phi}{\partial y^k}(y) \left(\frac{\partial f^k}{\partial x^i}(x) V^i(x) \right) = \langle d\phi, f_*V \rangle(f(x)) = \langle f^*d\phi, V \rangle(f(x)). \end{aligned}$$

Folglich gilt die Formel

$$f^*d\phi = d(f^*\phi),$$

d.h., *Differential und pull-back vertauschen*.

A.2.4 Tensorfelder

Tensorbündel An jedem Punkt P der Mannigfaltigkeit sitzt der Tangentialraum $T_P M$ und sein Dualraum $T_P^* M$. Damit können wir an jedem Punkt gemäß dem Vorgehen in Abschnitt 3.1.3 die Tensoralgebra über $T_P M$ aufbauen. Dem entsprechend gibt es dann für jeden Tensor-Typ (r, s) an jedem Punkt einen Vektorraum $T_P^{(r,s)} M$ von (r, s) -Tensoren. Wiederum kann man alle diese Vektorräume zu einem Bündel über M zusammenfassen, dem **Tensorbündel** $T^{(r,s)} M$ vom Typ (r, s) . Und wieder kann man Karten konstruieren, die es erlauben, diese Bündel als differenzierbare Mannigfaltigkeiten aufzufassen.

Tensorfeld Ein **Tensorfeld** vom Typ (r, s) ist in vollständiger Analogie dann eine differenzierbare Abbildung $M \rightarrow T_P^{(r,s)} M$, mit der Eigenschaft, dass jedem Punkt P ein Tensor auf dem Tangentialraum $T_P M$ in P zugeordnet wird, also ein Schnitt im Tensorbündel $T^{(r,s)} M$.

Mit Tensorfeldern lassen sich genau die gleichen Rechenoperationen durchführen wie mit Tensoren, da die Operationen punktweise definiert werden. So ist beispielsweise das äußere Produkt $T_1 \otimes T_2$ zweier Tensorfelder T_1 und T_2 dadurch definiert, dass man an jedem Punkt $P \in M$ das äußere Produkt $T_1(P) \otimes T_2(P)$ der beiden Werte bei P bildet (die ja beide Tensoren über dem gleichen Vektorraum $T_P M$ am Punkt P sind). Gleiches gilt sinngemäß für die anderen Tensoroperationen wie Addition, Skalarmultiplikation, Kontraktionen und Symmetrieoperationen. Zu beachten ist lediglich, dass man Tensorfelder nicht nur mit reellen Zahlen skalar multiplizieren kann, sondern sogar mit Skalarfeldern, also an jedem Punkt mit einer anderen Zahl.

Ebenso wie Vektor- und Kovektorfelder lassen sich auch Tensorfelder in lokalen Koordinaten angeben. Wählen wir eine Koordinatenbasis $\frac{\partial}{\partial x^i}$ mit zugehöriger dualer Basis dx^k in jedem Tangentialraum $T_P M$ in der Koordinatenumgebung, dann erhalten wir die Koordinaten eines (r, s) -Tensorfeldes V dadurch, dass wir das Tensorfeld (welches ja an jedem Punkt eine multilineare Abbildung $V(P)$ von r Kovektoren und s Vektoren ist) auf alle möglichen Kombinationen der Basisvektoren anwenden:

$$V^{i_1 \dots i_r}_{k_1 \dots k_s}(x[P]) = V(P)(dx^{i_1}, \dots, dx^{i_r}, \frac{\partial}{\partial x^{k_1}}, \dots, \frac{\partial}{\partial x^{k_s}}).$$

Die Komponenten des Tensorfeldes in Bezug auf eine Basis sind also n^{r+s} Funktionen der Koordinaten. Mithilfe dieser Komponentenfunktionen kann man das Tensorfeld vollständig rekonstruieren. Es gilt nämlich die Darstellung

$$V(P) = \sum_{\substack{i_1, \dots, i_r \\ k_1, \dots, k_s}} V^{i_1 \dots i_r}_{k_1 \dots k_s}(x[P]) \frac{\partial}{\partial x^{i_1}} \otimes \dots \otimes \frac{\partial}{\partial x^{i_r}} \otimes dx^{k_1} \otimes \dots \otimes dx^{k_s}.$$

Übung 1.8: Verifizieren Sie diese Formel.

Beispiele

- (i) Wie schon im Kap. 3 werden wir Vektorfelder mit Tensorfeldern vom Typ $(1, 0)$ und Kovektorfelder mit Tensorfeldern vom Typ $(0, 1)$ identifizieren.
- (ii) Ebenso werden wir Skalarfelder per Konvention mit Tensorfeldern vom Typ $(0, 0)$ gleichsetzen.
- (iii) Ein spezielles Tensorfeld ist der Kronecker-Tensor δ , der in jeder Karte konstante Koeffizienten hat, nämlich

$$\delta(P) = \sum_i dx^i \otimes \frac{\partial}{\partial x^i}.$$

- (iv) Ein $(0, 2)$ -Tensorfeld g hat in einer Karte die Form

$$g(P) = \sum_{i,k} g_{ik}(x[P]) dx^i \otimes dx^k.$$

- (v) Ein $(1, 3)$ -Tensorfeld R hat in einer Karte die Form

$$R(P) = \sum_{i,k,l,m} R^i_{klm}(x[P]) \frac{\partial}{\partial x^i} \otimes dx^k \otimes dx^l \otimes dx^m.$$

Wie schon im Kap. 3 werden wir im Folgenden ein (abstraktes) Tensorfeld mit Hilfe seiner Indexstruktur kennzeichnen, also indem wir beispielsweise mit v^a ein kontravariantes Vektorfeld oder mit ω_{ab} ein zweifach kovariantes Tensorfeld bezeichnen. Dabei sind nicht die Komponenten gemeint, d.h. wir betrachten die Indizes $a, b, \dots$ als abstrakte Symbole und nicht als Platzhalter für die natürlichen Zahlen.

Mit dieser Übereinkunft schreiben wir z.B. für den Kronecker-Tensor δ^a_b , für die Tensorfelder in Beispiel (iv) und (v) g_{ab} und R^a_{bcd} , etc. Damit lassen sich wie schon zuvor die Tensoroperationen leicht darstellen. Z.B. würde $R_{ab} = R^e_{aeb}$ das Tensorfeld $\tilde{R}$ bezeichnen, dessen Koordinatendarstellung durch

$$\tilde{R}(P) = \sum_{k,l} \left(\sum_m R^m_{kml}(x[P]) \right) dx^k \otimes dx^l$$

gegeben ist.

Bemerkung Pull-back bzw. Push-forward wurden oben für 1-Formen bzw. Vektorfelder eingeführt. Es ist ziemlich offensichtlich, dass man für beliebige Tensoren weder die eine noch die andere Abbildung definieren kann. Ist ein Tensorfeld rein kovariant, dann lässt sich der Pull-back definieren, ist es rein kontravariant, dann gibt es einen Push-forward. Handelt es sich bei der Abbildung f um einen Diffeomorphismus, dann kann man mithilfe von f^{-1} auch die jeweiligen Umkehrungen definieren und dann ist es auch möglich, beliebige Tensorfelder von M nach N zu verpflanzen.

A.3 Zusammenhang, kovariante Ableitung

Wie wir gesehen haben, ist die Differentiation eines Skalarfeldes eine wohldefinierte Operation. Das Differential eines Skalarfeldes, also eines $(0, 0)$ -Tensorfelds ist ein Kovektorfeld, ein $(0, 1)$ -Tensorfeld. In Hinblick auf unser Ziel, Analysis von Tensorfeldern zu betreiben, müssen wir uns jetzt fragen, ob und wie man diese Operation des Differenzierens auf allgemeine Tensorfelder verallgemeinern kann. Wir erwarten, dass sie ein (r, s) -Tensorfeld in ein $(r, s + 1)$ -Tensorfeld abbildet, denn es sollte mit jedem Tangentialvektor eine Richtungsableitung verbunden sein, die angibt, wie sich das betrachtete Tensorfeld in Richtung des Tangentialvektors ändert.

Es ist möglich, eine solche allgemeine Differentiationsoperation einzuführen, jedoch nur, indem man eine der Mannigfaltigkeit eine zusätzliche Struktur verleiht. Um zu sehen, was hier benötigt wird betrachten wir das einfachste nicht-triviale Beispiel, nämlich ein Vektorfeld V .

In einer Kartenumgebung U von M mit Koordinaten $(x^i)_{i=1:n}$ hat das Vektorfeld die Darstellung

$$\sum_i V^i(x) \frac{\partial}{\partial x^i}.$$

Nun könnten wir versucht sein, die Ableitung des Vektorfeldes einfach dadurch zu definieren, dass wir die Koordinatendarstellung der Ableitung durch die Komponenten

$$\frac{\partial V^i(x)}{\partial x^k}$$

definieren. Diese sehen so aus wie die Komponenten eines $(1, 1)$ -Tensorfelds. Um dies aber zu verifizieren, müssen wir nachweisen, dass sich diese Komponenten unter Kartenwechsel richtig transformieren. Es sei also U' eine weitere Kartenumgebung mit Koordinaten $(y^k)_{k=1:n}$, so dass V die Darstellung

$$\sum_k \tilde{V}^k(y) \frac{\partial}{\partial y^k}$$

in U' besitzt. Da ein Vektorfeld ein $(1, 0)$ -Tensorfeld ist, gilt

$$V^i(x(y)) = \sum_k \tilde{V}^k(y) \frac{\partial x^i}{\partial y^k}(y) = \sum_k \tilde{V}^k(y) J^i_k(y),$$

wobei wir hier zur Vereinfachung der Schreibweise die Jacobi-Matrix der Koordinatentransformation $x^i = x^i(y^k)$ eingeführt haben. Nun zeigt jedoch eine kurze Rechnung, dass

$$\begin{aligned} \frac{\partial V^i}{\partial x^k}(x(y)) &= \sum_{lm} \left[\frac{\partial \tilde{V}^l}{\partial y^m}(y) \frac{\partial y^m}{\partial x^k}(x) \frac{\partial x^i}{\partial y^l}(y) + \tilde{V}^l(y) \frac{\partial^2 x^i}{\partial y^l \partial y^m}(y) \frac{\partial y^m}{\partial x^k}(x(y)) \right] \\ &= \sum_{lm} \left[\frac{\partial \tilde{V}^l}{\partial y^m}(y) (J^{-1})^m_k(x(y)) J^i_l(y) + \tilde{V}^l(y) \frac{\partial J^i_l}{\partial y^m}(y) (J^{-1})^m_k(x(y)) \right]. \end{aligned}$$

Dies ist offensichtlich kein tensorielles Transformationsverhalten wie man am deutlichsten an dem Term mit den zweiten Ableitungen der Koordinatentransformation erkennt. Tensorielles Transformationsverhalten bedeutet insbesondere, dass sich die Komponenten *homogen* transformieren, was hier nicht der Fall ist.

Wollen wir tensorielles Transformationsverhalten erreichen, dann müssen wir die Komponenten der Ableitung des Vektorfeldes 'korrigieren', d.h. so abändern, dass der inhomogene Term bei Kartenwechsel wegfällt. Da dieser Term linear in den Feldkomponenten ist, machen wir für die Komponenten der Ableitung von V den Ansatz

$$\frac{\partial V^i(x)}{\partial x^k} + \sum_l \Gamma_{kl}^i(x) V^l(x), \quad (\text{A.3.1})$$

d.h., wir nehmen an, dass uns n^3 Größen $\Gamma_{kl}^i(x)$ bezüglich der Koordinaten (x^i) gegeben sind, die wir zur Korrektur des falschen Transformationsverhaltens verwenden. Dies kann aber nur dann funktionieren, wenn diese Größen selbst ein ganz bestimmtes Transformationsverhalten aufweisen. Dieses darf natürlich ebenfalls nicht tensoriell sein, denn sonst wäre die korrigierte Ableitung wieder kein Tensor. Vielmehr muss das nicht-tensorielle Verhalten der partiellen Ableitung vom nicht-tensoriellen Transformationsverhalten der Γ aufgehoben werden. Um dieses Verhalten zu bestimmen, betrachten wir wieder einen Kartenwechsel:

$$\begin{aligned} & \frac{\partial V^i}{\partial x^k}(x(y)) + \sum_l \Gamma_{kl}^i(x(y)) V^l(x(y)) \\ &= \sum_{lm} \left[\frac{\partial \tilde{V}^l}{\partial y^m}(y) (J^{-1})^m_k(x(y)) J^l_l(y) + \tilde{V}^l(y) \frac{\partial J^l_l}{\partial y^m}(y) (J^{-1})^m_k(x(y)) \right. \\ & \quad \left. + \Gamma_{km}^i(x(y)) \tilde{V}^l(y) J^m_l(y) \right] \\ &= \sum_{lm} \left[\frac{\partial \tilde{V}^l}{\partial y^m}(y) + \sum_n \tilde{\Gamma}_{mn}^l(y) \tilde{V}^n(y) \right] (J^{-1})^m_k(x(y)) J^l_l(y) \\ &+ \sum_{lm} \tilde{V}^l(y) \left(\frac{\partial J^l_l}{\partial y^m}(y) (J^{-1})^m_k(x(y)) + \Gamma_{km}^i(x(y)) J^m_l(y) \right. \\ & \quad \left. - \sum_n \tilde{\Gamma}_{nl}^m(y) (J^{-1})^n_k(x(y)) J^i_m(y) \right). \end{aligned}$$

Wenn also die Ableitung des Vektorfeldes wieder ein Tensorfeld sein soll, dann muss die letzte Summe für beliebige $\tilde{V}^l$ verschwinden. Aus dieser Bedingung ergibt sich das Transformationsverhalten der Γ_{kl}^i bei Kartenwechsel:

$$\Gamma_{kl}^i(x(y)) = \sum_{jm} \left(\sum_n \left[J^i_n(y) \tilde{\Gamma}_{mj}^n(y) \right] - \frac{\partial J^i_j}{\partial y^m}(y) \right) (J^{-1})^m_k(x(y)) (J^{-1})^j_l(x(y)).$$

Wir kommen also nun zu dem folgenden Schluss: wenn wir für jede Karte n^3 Größen $\Gamma_{kl}^i(x)$ als Funktionen der Koordinaten (x^i) vorgeben, die sich bei Kartenwechsel wie oben angegeben transformieren, dann ist es möglich, zu jedem Vektorfeld eine Ableitung zu berechnen und zwar so, dass diese wieder ein Tensorfeld ist. Aus diesem Grund nennt man diese Operation **kovariante Ableitung**. Die Größen $\Gamma_{kl}^i(x)$ nennt man **Zusammenhangskoeffizienten**. Sie definieren eine invariante Struktur auf der Mannigfaltigkeit, nämlich einen sogenannten **Zusammenhang** oder auch Konnexion. Ein Zusammenhang ist *kein Tensorfeld* auf M , trotzdem ist er invariant, also koordinatunabhängig, definiert.

Die kovariante Ableitung ordnet jedem Vektorfeld V ein $(1, 1)$ -Tensorfeld zu. Für jeden Tangentialvektor $t \in T_P M$ in einem Punkt P erhält man daraus einen weiteren Tangentialvektor, dessen Komponenten bzgl. einer Karte durch

$$\sum_k t^k \left(\frac{\partial V^i}{\partial x^k} + \sum_l \Gamma_{kl}^i V^l \right)$$

gegeben sind. Dieser Vektor ist die kovariante Ableitung von V in Richtung t . Dieser Begriff verallgemeinert die Richtungsableitung von Skalarfeldern auf Vektorfelder.

Vielfach wird ein Zusammenhang auf M mit einem **Ableitungsoperator** ∇ gleichgesetzt. Dabei handelt es sich um eine Abbildung, die jedem Vektorfeld V ein $(1, 1)$ -Tensorfeld ∇V zuordnet und dabei folgenden Forderungen genügt:

1. $\nabla : V \mapsto \nabla V$ ist additiv: für je zwei Vektorfelder U und V gilt

$$\nabla(U + V) = \nabla U + \nabla V.$$

2. ∇ erfüllt die Leibniz-Regel (Produktregel): für jedes Vektorfeld V und jedes Skalarfeld ϕ gilt

$$\nabla(\phi V) = d\phi \otimes V + \phi \nabla V.$$

Man prüft leicht nach, dass diese Bedingungen für die oben definierte kovariante Ableitung erfüllt sind. Wir können also für die kovariante Ableitung eines Vektorfeldes V nun auch ∇V schreiben und für die kovariante Richtungsableitung in Richtung t schreiben wir auch $\nabla_t V$. Die Wirkung des Ableitungsoperators auf Skalarfelder wird durch das Differential gegeben, d.h. man legt fest, dass für beliebige Skalarfelder ϕ gilt

$$\nabla \phi = d\phi.$$

Wie sind die Zusammenhangskoeffizienten durch den Ableitungsoperator gegeben? Dazu betrachten wir die Koordinatenbasis in einer Kartenumgebung. An jedem Punkt P mit Koordinaten $x[P]$ können wir die Richtungsableitung des Koordinatenvektors $\frac{\partial}{\partial x^k}$ in Richtung eines anderen Koordinatenvektors $\frac{\partial}{\partial x^i}$ bestimmen. Dieses ist wieder

ein Tangentialvektor in P , also durch die Basisvektoren linear kombinierbar. Es muß also eine Relation der Form

$$\nabla_{\frac{\partial}{\partial x^i}} \frac{\partial}{\partial x^k} = \sum_l \Gamma_{ik}^l \frac{\partial}{\partial x^l} \quad (\text{A.3.2})$$

bestehen. Diese Entwicklungskoeffizienten sind gerade die Zusammenhangskoeffizienten. Dies wird noch einleuchtender, wenn man die kovariante Ableitung eines Vektorfeldes unter Benutzung von ∇ berechnet. Das Vektorfeld V habe die Koordinatendarstellung

$$V = \sum_i V^i \partial_i,$$

wobei wir $\partial_i = \frac{\partial}{\partial x^i}$ gesetzt haben. Nun ist nach den Regeln für ∇

$$\nabla V = \sum_i dV^i \otimes \partial_i + V^i \nabla(\partial_i).$$

Für die Richtungsableitung $\nabla_{\partial_k} V$ entlang des Basisvektors ∂_k finden wir dann

$$\begin{aligned} \nabla_{\partial_k} V &= \sum_i \langle dV^i, \partial_k \rangle \partial_i + \sum_i V^i \nabla_{\partial_k}(\partial_i) = \sum_i \frac{\partial V^i}{\partial x^k} \partial_i + \sum_{il} V^i \Gamma_{ki}^l \partial_l \\ &= \sum_i \left(\frac{\partial V^i}{\partial x^k} + \sum_l V^l \Gamma_{kl}^i \right) \partial_i. \end{aligned}$$

Man vergleiche dies mit (A.3.1).

Übung 1.9: Zeigen Sie, wie sich auch das Transformationsverhalten der Zusammenhangskoeffizienten aus (A.3.2) herleiten lässt.

In unserer Indexschreibweise wird der Ableitungsoperator mit einem (abstrakten) Index versehen, der andeutet, dass sich durch den Prozess des Ableitens der kovariante Grad um eins erhöht. So ist ja die Ableitung eines Skalarfeldes ein Kovektorfeld und wir schreiben dementsprechend

$$\nabla_a \phi$$

für die kovariante Ableitung von ϕ . Analog wird aus einem Vektorfeld V^b durch kovariante Ableitung ein $(1, 1)$ -Tensorfeld $\nabla_a V^b$. Die Richtungsableitung in Richtung eines Vektors t^a ist dann gegeben durch $t^a \nabla_a V^b = (\nabla_t V)^b$. Mit dieser symbolischen Bezeichnung wird es einfacher, die kovariante Ableitung für beliebige Tensorfelder zu definieren. Dies geschieht dadurch, dass man Regeln aufstellt, nach denen man die Ableitung berechnen will. Diese sollten möglichst natürlich sein und dem Wesen einer Ableitung entsprechen. Im einzelnen fordert man für die kovariante Ableitung (bzw. den Ableitungsoperator)

1. Linearität: für je zwei Tensorfelder $T^{b\dots}_{c\dots}$ und $S^{b\dots}_{c\dots}$ gleichen Typs ist

$$\nabla_a \left(T^{b\dots}_{c\dots} + S^{b\dots}_{c\dots} \right) = \nabla_a T^{b\dots}_{c\dots} + \nabla_a S^{b\dots}_{c\dots}.$$

2. Produktregel: für je zwei Tensorfelder $T^{b\dots c\dots}$ und $S^{d\dots e\dots}$ beliebigen Typs ist

$$\nabla_a \left(T^{b\dots c\dots} S^{d\dots e\dots} \right) = \left(\nabla_a T^{b\dots c\dots} \right) S^{d\dots e\dots} + T^{b\dots c\dots} \left(\nabla_a S^{d\dots e\dots} \right).$$

3. Vertauschen mit Kontraktionen: die kovariante Ableitung eines kontrahierten Tensorfeldes ist das gleiche wie die Kontraktion (über entsprechende Indizes) der kovarianten Ableitung des Tensorfeldes.

Einige Bemerkungen hierzu:

- Der Kronecker-Tensor hat in jedem Koordinatensystem die gleiche Gestalt. Er ist an jedem Punkt der Koordinatenumgebung gleich und damit auf ganz M . Für jeden sinnvollen Ableitungsbegriff sollte seine Ableitung daher verschwinden. In der Tat folgt dies einfach aus der Relation $\delta_c^b = \delta_d^b \delta_c^d$ durch Differenzieren:

$$\nabla_a \delta_c^b = \nabla_a \delta_d^b \delta_c^d + \delta_d^b \nabla_a \delta_c^d = 2 \nabla_a \delta_c^b.$$

Dabei wurde die Produktregel und die Vertauschung mit Kontraktionen benutzt.

- Mit diesem Resultat lässt sich auch die Vertauschung mit Kontraktionen etwas klarer formulieren. Es sei $T^{a\dots b\dots c\dots d\dots}$ ein beliebiges Tensorfeld und $T^{a\dots b\dots c\dots b\dots} = \delta_b^d T^{a\dots b\dots c\dots d\dots}$ eine beliebige Kontraktion. Dann gilt

$$\nabla_e \left(\delta_b^d T^{a\dots b\dots c\dots d\dots} \right) = \delta_b^d \left(\nabla_e T^{a\dots b\dots c\dots d\dots} \right).$$

Die Konstanz des Kronecker-Tensors ist also äquivalent zu Regel 3.

- Regel 2. ist die Verallgemeinerung der oben formulierten Leibnizregel.
- Mit diesen Regeln ist die kovariante Ableitung eindeutig festgelegt, wenn man weiss, wie sie auf Vektorfelder wirkt, d.h. wenn man für eine (und damit für jede) Karte die Zusammenhangskoeffizienten kennt. Dies folgt aus der unten hergeleiteten Wirkung auf Kovektorfelder und aus der Tatsache, dass jedes Tensorfeld in einer Kartenumgebung als eine Linearkombination von äußeren Produkten der Basis(ko)vektoren darstellbar ist.

Die Wirkung der kovarianten Ableitung auf ein Kovektorfeld ω_a ergibt sich aus der folgenden Betrachtung. Für jedes Vektorfeld V^a ist

$$\nabla_e (\omega_a V^a) = \nabla_e \omega_a V^a + \omega_a \nabla_e V^a.$$

Setzen wir hier den Spezialfall ein, wo V^a ein Koordinatenbasisvektor ∂_k^a und ω_a ein Koordinatendifferential dx_a^j ist⁵, so gilt (man beachte, dass $\langle dx^j, \partial_k \rangle = dx_a^j \partial_k^a = \delta_k^j$ ist)

$$0 = \nabla_e \delta_k^j = \left(\nabla_e dx_a^j \right) \partial_k^a + dx_a^j \left(\nabla_e \partial_k^a \right).$$

⁵Man beachte hier die unterschiedliche Qualität der Indizes: 'k' und 'j' stehen für beliebige natürliche Zahlen zwischen 1 und n während 'a' hier als (zusätzlicher) Indikator dafür steht, dass es sich um Kovektoren bzw. Vektoren handelt.

Überschieben wir diese Gleichung jetzt mit ∂_l^e , d.h. berechnen wir die Richtungsableitung in Richtung des Basisvektors ∂_l^e , dann folgt

$$\begin{aligned} 0 &= \left(\nabla_e dx_a^j \right) \partial_l^e \partial_k^a + dx_a^j \left(\nabla_{\partial_l} \partial_k^a \right) \\ &= \left(\nabla_e dx_a^j \right) \partial_l^e \partial_k^a + dx_a^j \left(\sum_i \Gamma_{lk}^i(x) \partial_l^i \right) = \left(\nabla_e dx_a^j \right) \partial_l^e \partial_k^a + \Gamma_{lk}^j(x) \end{aligned}$$

Das $(0,2)$ -Tensorfeld $\nabla_e dx_a^j$ hat die Koordinatendarstellung

$$\nabla_e dx_a^j = \sum_{im} \alpha_{im}^j(x) dx_e^i dx_a^m$$

mit zu bestimmenden Koeffizientenfunktionen $\alpha_{kl}^i(x)$. Setzen wir dies in die obige Gleichung ein, so folgt

$$\begin{aligned} -\Gamma_{lk}^j &= \left(\nabla_e dx_a^j \right) \partial_l^e \partial_k^a = \sum_{im} \alpha_{im}^j(x) \left(dx_e^i dx_a^m \right) \partial_l^e \partial_k^a \\ &= \sum_{im} \alpha_{im}^j(x) \left(dx_e^i \partial_l^e \right) \left(dx_a^m \partial_k^a \right) = \alpha_{lk}^j(x). \end{aligned}$$

Damit haben wir die beiden wesentlichen Relationen für die Zusammenhangskoeffizienten gefunden:

$$\nabla_e \partial_i^a = \sum_{lk} \Gamma_{li}^k(x) dx_e^l \partial_k^a, \quad \nabla_e dx_a^i = - \sum_{lk} \Gamma_{lk}^i(x) dx_e^l dx_a^k. \quad (\text{A.3.3})$$

Hiermit kann man die kovariante Ableitung für beliebige Tensorfelder bestimmen.

Als Beispiel sei hier ein $(0,2)$ -Tensorfeld betrachtet, das die Koordinatendarstellung

$$g_{ab} = \sum_{ik} g_{ik} dx_a^i dx_b^k$$

besitzt (das Argument x wird unterdrückt). Die kovariante Ableitung berechnet man mithilfe der Rechenregeln für ∇_e wie folgt

$$\begin{aligned} \nabla_e g_{ab} &= \sum_{ik} \left(\nabla_e g_{ik} dx_a^i dx_b^k + g_{ik} \nabla_e dx_a^i dx_b^k + g_{ik} dx_a^i \nabla_e dx_b^k \right) \\ &= \sum_{ik} \left(\sum_l \frac{\partial g_{ik}}{\partial x^l} dx_e^l dx_a^i dx_b^k - g_{ik} \sum_{lm} \Gamma_{lm}^i dx_e^l dx_a^m dx_b^k \right. \\ &\quad \left. - g_{ik} \sum_{lm} \Gamma_{lm}^k dx_a^i dx_e^l dx_b^m \right). \end{aligned}$$

Eine leichte Umordnung und Umbenennung der Summationsindizes ergibt schließlich

$$\nabla_e g_{ab} = \sum_{ikl} \left(\frac{\partial g_{ik}}{\partial x^l} - \sum_m g_{mk} \Gamma_{li}^m - \sum_m g_{lm} \Gamma_{ik}^m \right) dx_e^l dx_a^i dx_b^k. \quad (\text{A.3.4})$$

Das bedeutet, die Komponenten der kovarianten Ableitung von g_{ab} in einer Basis ist durch den Ausdruck in Klammern gegeben. Auf ähnliche Weise berechnet man die kovariante Ableitung für Tensorfelder anderen Typs.

Die Bezeichnung 'Zusammenhang' kommt daher, dass wir mit der kovarianten Ableitung ein Werkzeug in der Hand haben, mit dem wir Vektoren *an verschiedenen Punkten* der Mannigfaltigkeit miteinander in Beziehung bringen können und damit einen Zusammenhang zwischen ihnen herstellen. Um dies zu sehen, betrachten wir zwei Punkte P und Q in M und eine glatte Kurve $\gamma(\lambda)$ zwischen ihnen, so dass $P = \gamma(0)$ und $Q = \gamma(1)$ gilt. An jedem Punkt der Kurve gibt es den Tangentialvektor $t = \dot{\gamma}(\lambda)$. Es sei nun ein Vektorfeld V entlang der Kurve gegeben, dann nennen wir V ein **paralleles Vektorfeld** entlang γ , falls an jedem Punkt der Kurve

$$\nabla_t V = 0$$

ist. Man sagt auch, der Vektor $V(P)$ werde entlang γ parallel transportiert zum Punkt Q . Dieser **Paralleltransport** stellt die Verbindung zwischen den Tangentialräumen an verschiedenen Punkten der Mannigfaltigkeit her.

Offensichtlich ist der Paralleltransport für Tensoren vollkommen analog definiert. Der Paralleltransport ist eine Verallgemeinerung der Parallelverschiebung der euklidischen Geometrie auf beliebige Mannigfaltigkeiten. Im Falle des euklidischen Raums kann man den Ableitungsoperator dadurch festlegen, dass man die Zusammenhangskoeffizienten *bzgl. der kartesischen Koordinaten* Null setzt. Dies bedeutet nicht, dass die Zusammenhangskoeffizienten in jedem Koordinatensystem verschwinden. Vielmehr kann man aus der Koordinatentransformation die entsprechenden Zusammenhangskoeffizienten berechnen und damit den Ableitungsoperator in beliebigen Koordinaten ausdrücken.

In der euklidischen Geometrie sind Geraden dadurch ausgezeichnet, dass sie bei Parallelverschiebung entlang sich selbst in sich übergehen. Anders ausgedrückt, verschieben wir den Tangentialvektor an die Gerade in Richtung der Geraden, also parallel zu sich, so ändert sich seine Richtung nicht. In Analogie dazu betrachten wir im allgemeinen Fall eine Kurve auf einer Mannigfaltigkeit und in einem beliebigen ihrer Punkte, P , ihren Tangentialvektor t . Gilt in P

$$\nabla_t t \sim t$$

so nennen wir die Kurve *gerade in P* . Gilt die Gleichung an jedem ihrer Punkte, dann ist die Kurve eine **Autoparallele**

Die so definierte kovariante Ableitung gestattet es uns nun, Tensoren an verschiedenen Punkten einer Mannigfaltigkeit miteinander in Beziehung zu setzen. Insbesondere können wir jetzt die Änderung eines Tensorfeldes in einem beliebigen Punkt in einer beliebigen Richtung berechnen. Was wir dazu benötigen sind die n^3 Zusammenhangskoeffizienten in einer Koordinatenumgebung. Nun sind dies offensichtlich viele Freiheiten, die man bei der Wahl dieser Koeffizienten hat und es ist nicht klar, wie man diese wählen soll. Es ist beispielsweise immer möglich, die Koeffizienten so zu wählen,

dass ein bestimmtes vorgegebenes Vektorfeld kovariant konstant ist, also eine identisch verschwindende kovariante Ableitung besitzt.

Um etwas über die Freiheiten zu erfahren und damit auch über die Möglichkeiten, sie einzuschränken nehmen wir an, dass es zusätzlich zu ∇_a einen weiteren Ableitungsoperator $\tilde{\nabla}_a$ auf M gibt. Wir interessieren uns für den 'Unterschied' zwischen den beiden Operatoren. Offensichtlich gilt für jedes Skalarfeld $\phi : M \rightarrow \mathbb{R}$

$$\nabla\phi = d\phi = \tilde{\nabla}\phi,$$

d.h., die Wirkung der beiden Ableitungsoperatoren auf Skalarfelder ist gleich.

Dagegen ergibt sich für Vektorfelder im allgemeinen ein Unterschied, denn in einer Koordinatenumgebung gilt

$$\nabla_{\frac{\partial}{\partial x^i}} \frac{\partial}{\partial x^k} = \sum_l \Gamma_{ik}^l \frac{\partial}{\partial x^l}, \quad \tilde{\nabla}_{\frac{\partial}{\partial x^i}} \frac{\partial}{\partial x^k} = \sum_l \tilde{\Gamma}_{ik}^l \frac{\partial}{\partial x^l}.$$

Betrachten wir die Differenz der Zusammenhangskoeffizienten

$$C^l_{ik} = \tilde{\Gamma}_{ik}^l - \Gamma_{ik}^l.$$

Berechnen wir diese Koeffizienten nach einer Koordinatentransformation, so stellt sich heraus, dass sie sich transformieren wie die Komponenten eines $(1, 2)$ -Tensors.

Übung 1.10: Beweisen Sie diese Behauptung.

Dies liegt letztlich daran, dass der inhomogene Term in der Transformation der Zusammenhangskomponenten nicht von den Γ 's abhängt, sondern nur von der Koordinatentransformation, und diese ist für beide Zusammenhangskoeffizienten die gleiche. Wir können also in der Koordinatenumgebung den Tensor

$$\sum_{ikl} [\tilde{\Gamma}_{ik}^l - \Gamma_{ik}^l] \partial_l \otimes dx^i \otimes dx^k$$

anschreiben und wissen, dass dieser Tensor überall definiert ist. Damit ist gezeigt: *die Differenz zweier Ableitungsoperatoren ist ein $(1, 2)$ -Tensorfeld*. Es gibt also ebenso viele Ableitungsoperatoren wie es $(1, 2)$ -Tensorenfelder gibt. In unserer Indexschreibweise haben wir

$$\tilde{\nabla}_a V^b = \nabla_a V^b + C_a{}^b{}_c V^c, \quad \tilde{\nabla}_a \omega_b = \nabla_a \omega_b - C_a{}^c{}_b \omega_c$$

für beliebige Vektorfelder V^b bzw. Kovektorfelder ω_b . Es ist offensichtlich, dass mit gegebenem Ableitungsoperator ∇ jedes Tensorfeld $C_a{}^b{}_c$ auf diese Weise einen weiteren Ableitungsoperator definiert, denn die beiden Kriterien für Ableitungsoperatoren sind erfüllt.

Um diese riesige Auswahl an Ableitungsoperatoren einschränken zu können, betrachten wir zunächst die **Torsion** des Ableitungsoperators. Dies ist das $(1, 2)$ -Tensorfeld Torsion

T^a_{bc} , welches durch die Relation

$$(\nabla_a \nabla_b - \nabla_b \nabla_a) \phi = T^c_{ab} \nabla_c \phi$$

für beliebige Skalarfelder ϕ erklärt wird. Die Torsion ist also ein Maß dafür, wie weit die kovariante Ableitung angewandt auf Skalarfelder kommutativ ist. Kann man Ableitungsoperatoren finden, deren Torsion verschwindet? Um diese Frage zu beantworten, betrachten wir die zwei Ableitungsoperatoren ∇ und $\tilde{\nabla}$ und ihre jeweilige Torsion. Es gilt nun

$$\tilde{\nabla}_a \tilde{\nabla}_b \phi = \tilde{\nabla}_a (\nabla_b \phi) = \nabla_a \nabla_b \phi - C^c_{ab} \nabla_c \phi$$

und damit

$$\tilde{T}^c_{ab} = T^c_{ab} - C^c_{ab} + C^c_{ba} = T^c_{ab} - 2C^c_{[ab]}.$$

Sei nun $\tilde{\nabla}$ torsionsfrei, also $\tilde{T}^c_{ab} = 0$. Dann erfüllt der Differenztensor die Gleichung

$$2C^c_{[ab]} = T^c_{ab}.$$

Dies ist eine lineare Gleichung für C^c_{ab} . Eine partikuläre Lösung dieser Gleichung ist einfach

$$C^c_{ab} = T^c_{ab}.$$

Die homogene Gleichung

$$C^c_{[ab]} = 0$$

besagt, dass jede ihrer Lösungen symmetrisch in den unteren Indizes ist, so dass wir die allgemeine Lösung in der Form

$$C^c_{ab} = T^c_{ab} + S^c_{ab}$$

schreiben können, wobei $S^c_{ab} = S^c_{ba}$ sein muss, aber sonst beliebig gewählt werden darf.

Wir erhalten durch diese Rechnung folgendes Ergebnis

Satz A.1. *Auf einer Mannigfaltigkeit kann man immer einen torsionsfreien Zusammenhang einführen. Je zwei torsionsfreie Zusammenhänge unterscheiden sich durch ein Tensorfeld C^c_{ab} mit $C^c_{ab} = C^c_{ba}$.*

Wir werden ab jetzt immer annehmen, dass ein Ableitungsoperator torsionsfrei ist und dass dem entsprechend der Differenztensor symmetrisch in den unteren Indizes ist.

A.4 Der Riemann-Tensor, Krümmung

Wie stellt man fest ob eine Fläche und, allgemeiner, eine Mannigfaltigkeit gekrümmt ist? Zunächst müssen wir zwei Arten von Krümmung unterscheiden. Aus unserer Anschauung entnehmen wir einen Begriff von Krümmung der darauf beruht, wie sich

eine Fläche im drei-dimensionalen Raum verbiegt. Dies führt auf den Begriff der 'äußeren Krümmung', der mit Hilfe eines umgebenden Raumes definiert werden muss. Wir sind hier jedoch an der Krümmung der Raumzeit interessiert, die wir ohne Bezugnahme auf einen umgebenden Raum feststellen wollen. Wir benötigen daher den Begriff der 'inneren Krümmung' oder einfach 'Krümmung'.

Diese Krümmung beruht auf dem Begriff des Paralleltransports bzw. dessen infinitesimaler Version, der kovarianten Ableitung. Aus der Anschauung entnehmen wir, dass sich ein Vektor, der in der Ebene parallel entlang einer geschlossenen Kurve transportiert wird, wieder auf sich selbst abgebildet wird (vgl. Abb. A.5). Hingegen ist dies auf einer Kugel nicht so. Eine weitere Möglichkeit die Krümmung einer Mannigfaltigkeit

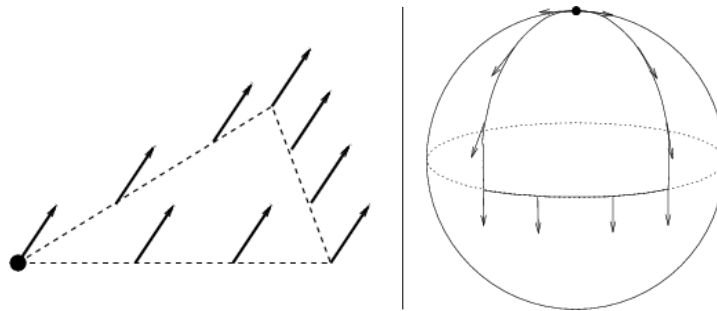


Abbildung A.5: Paralleltransport eines Vektors um eine geschlossene Kurve in der Ebene und auf der Kugeloberfläche

zu beurteilen besteht darin, zwei parallele 'Geraden' zu verfolgen und zu sehen, ob sie sich voneinander entfernen oder aber aufeinander zulaufen. In der Ebene oder im drei-dimensionalen euklidischen Raum bleiben parallele Geraden immer parallel, während dies im allgemeinen nicht so sein muss. Der Grad der Abweichung von der Parallelität ist ein Maß für die Krümmung.

Beide Möglichkeiten beruhen auf dem Begriff des Paralleltransports. Im ersten Fall ist das offensichtlich. Im zweiten Fall steckt der Paralleltransport im Begriff der geraden Linie als derjenigen Linie, deren Tangentialvektor parallel entlang sich selbst transportiert wird.

Wir kommen auf die erste Möglichkeit zurück (vgl. Abb. A.6) und betrachten einen Vektor, der ausgehend von P entlang der angegebenen Kurve parallel transportiert wird. Wir nehmen an, dass die Kurvenparameter ein Intervall Δt bzw. Δs durchlaufen. Uns interessiert die Differenz der Vektoren bei P vor und nach dem Transport. Wir sehen aus dem Bild, dass dabei zwei Paralleltransporte wesentlich sind, zuerst entlang der Kurve mit Parameter t nach rechts, wo aus dem Vektor 1 bei P der Vektor 2 entsteht, und dann entlang der Kurve mit Parameter s nach oben, wo sich Vektor 3 ergibt. Anschließend geht es in umgekehrter Reihenfolge wieder nach P zurück. Stellen wir uns jetzt vor, dass dieses ganze Bild infinitesimal zu sehen ist, dann ist es nicht verwunderlich, dass bei der Berechnung der Differenz der Vektoren 1 und 5 bei P der Kommutator

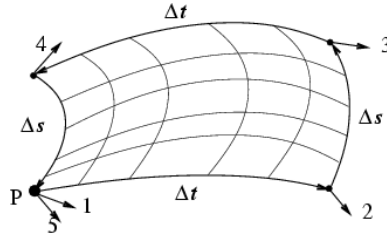


Abbildung A.6: Paralleltransport um eine geschlossene Kurve

der kovarianten Ableitung eine Rolle spielen muss. Wie dies im Detail zustande kommt, wird ausführlich im Buch von Wald [22] gezeigt.

Wir werden also darauf geführt, den folgenden Ausdruck zu betrachten

$$X_{ab}{}^c := \nabla_a \nabla_b V^c - \nabla_b \nabla_a V^c$$

für ein beliebiges Vektorfeld V^c . Offensichtlich ist dies ein $(1, 2)$ -Tensorfeld. Es ist nun eine erstaunliche Tatsache, dass dieser Ausdruck, der scheinbar von der zweiten Ableitung des Vektorfeldes abhängt, tatsächlich nur von dem Vektorfeld selbst abhängt. Es kommt nicht einmal die erste Ableitung vor. Wir zeigen dies folgendermaßen. Für ein beliebiges Skalarfeld f betrachten wir die zweifache kovariante Ableitung $\nabla_a \nabla_b (fV^c)$ und erhalten mit den üblichen Rechenregeln für die kovariante Ableitung

$$\begin{aligned} \nabla_a \nabla_b (fV^c) &= \nabla_a (\nabla_b f V^c + f \nabla_b V^c) \\ &= (\nabla_a \nabla_b f) V^c + \nabla_a f \nabla_b V^c + \nabla_b f \nabla_a V^c + f (\nabla_a \nabla_b V^c). \end{aligned}$$

Bei der Berechnung des Kommutators heben sich die ersten drei Terme weg (man beachte, dass wir uns auf torsionsfreie Zusammenhänge beschränken), so dass wir die einfache Relation

$$(\nabla_a \nabla_b - \nabla_b \nabla_a) (fV^c) = f (\nabla_a \nabla_b - \nabla_b \nabla_a) V^c$$

erhalten. Diese Eigenschaft ist ausschlaggebend dafür, dass der Wert des Tensorfeldes $X_{ab}{}^c$ an einem Punkt $P \in M$ nur vom Wert des Vektorfeldes V^c bei P abhängt. Denn nehmen wir an, dass V^c und $\tilde{V}^c$ zwei Vektorfelder sind, die bei P übereinstimmen. Dann ist offensichtlich $V^c(P) - \tilde{V}^c(P) = 0$. Nun kann man immer n glatte Funktionen $f^{(\alpha)}$ mit $f^{(\alpha)}(P) = 0$ und n Vektorfelder $X_{(\alpha)}^c$ finden, so dass

$$V^c - \tilde{V}^c = \sum_{\alpha=1}^n f^{(\alpha)} X_{(\alpha)}^c$$

gilt. Damit wird

$$\begin{aligned} X_{ab}{}^c - \tilde{X}_{ab}{}^c &= (\nabla_a \nabla_b - \nabla_b \nabla_a) (V^c - \tilde{V}^c) = (\nabla_a \nabla_b - \nabla_b \nabla_a) \left(\sum_{\alpha=1}^n f^{(\alpha)} X_{(\alpha)}^c \right) \\ &= \sum_{\alpha=1}^n f^{(\alpha)} (\nabla_a \nabla_b - \nabla_b \nabla_a) X_{(\alpha)}^c \end{aligned}$$

und somit offensichtlich

$$X_{ab}{}^c(P) - \tilde{X}_{ab}{}^c(P) = 0.$$

Das bedeutet, dass tatsächlich der Wert von $X_{ab}{}^c(P)$ nur vom Wert $V^c(P)$ abhängt. Anders ausgedrückt: die Abbildung

$$(T^a, S^b, V^c, \omega_d) \mapsto T^a S^b \omega_c [(\nabla_a \nabla_b - \nabla_b \nabla_a) V^c]$$

ist an jedem Punkt P eine lineare Abbildung nach $\mathbb{R}$, und daher ein $(1, 3)$ -Tensorfeld, das wir mit $R_{abc}{}^d$ bezeichnen. Das dadurch definierte Tensorfeld ist der **Riemann-Tensor**. Wir haben also die folgende Relation⁶

Riemann-Tensor

$$2\nabla_{[a} \nabla_{b]} V^d = (\nabla_a \nabla_b - \nabla_b \nabla_a) V^d = R_{abc}{}^d V^c.$$

Übung 1.11: Zeigen Sie, dass die entsprechende duale Version wie folgt lautet

$$2\nabla_{[a} \nabla_{b]} \omega_c = (\nabla_a \nabla_b - \nabla_b \nabla_a) \omega_c = -R_{abc}{}^d \omega_d.$$

Diese beiden Relationen, die es gestatten zweite Ableitungen durch Krümmungsterme auszudrücken, findet man auch unter dem Namen **Ricci-Identitäten**.

Ricci-Identitäten

Der Riemann-Tensor ist ein besonderer $(1, 3)$ -Tensor, der einige zusätzliche Eigenschaften besitzt. Da ist zuerst einmal die offensichtliche *Antisymmetrie* in seinen ersten beiden Indizes:

$$R_{abc}{}^d = -R_{bac}{}^d.$$

Außerdem gilt die **zyklische Identität**

zyklische Identität

$$R_{[abc]}{}^d = 0,$$

sowie die **Bianchi-Identität**

Bianchi-Identität

$$\nabla_{[e} R_{ab]c}{}^d = 0.$$

Als Beispiel beweisen wir die zyklische Identität, die Bianchi-Identität wird ganz ähnlich bewiesen. Wir betrachten zunächst einen beliebigen total antisymmetrischen $(0, 3)$ -Tensor T_{abc} . Es gilt

$$T_{abc} = T_{[abc]} = T_{[ab]c} = T_{a[bc]}.$$

⁶Das Vorzeichen vor dem Riemann-Tensor und seine Indexstellung unterliegen verschiedenen Konventionen, die leider nicht einheitlich sind.

Das erste Gleichheitszeichen ist die Definition der totalen Antisymmetrie. Wir bestimmen

$$T_{[[ab]c]} = \frac{1}{2} (T_{[abc]} - T_{[bac]}) = T_{[abc]}$$

und analog

$$T_{[a[bc]]} = \frac{1}{2} (T_{[abc]} - T_{[acb]}) = T_{[abc]}.$$

Wie wenden dies auf die dreifache total antisymmetrisierte Ableitung eines Skalarfeldes f an

$$\nabla_{[a} \nabla_b \nabla_{c]} f = \nabla_{[[a} \nabla_b] \nabla_{c]} f = \frac{1}{2} R_{[abc]}^d \nabla_d f$$

und

$$\nabla_{[a} \nabla_b \nabla_{c]} f = \nabla_{[a} \nabla_{[b} \nabla_{c]} f = 0,$$

da der Zusammenhang torsionsfrei ist.

Übung 1.12: Beweisen Sie die Bianchi-Identität.

Wie kann man nun den Riemann-Tensor eines gegebenen Ableitungsoperators konkret berechnen? Sei uns also ein Ableitungsoperator bzgl. eines Koordinatensystems durch seinen Zusammenhangskoeffizienten gegeben. Es gilt dann also

$$\nabla_{\partial_i} \partial_k = \sum_l \Gamma_{ik}^l \partial_l. \quad (\text{A.4.1})$$

Um die Koeffizienten des Riemann-Tensors zu berechnen, muss man wie üblich seine Wirkung auf den Basisvektoren bestimmen, also die Funktionen

$$R_{abc}{}^d \partial_i^a \partial_k^b \partial_l^c dx_d^m$$

berechnen. Dazu betrachten wir zunächst den allgemeinen Fall

$$X^a Y^b Z^c R_{abc}{}^d = X^a Y^b (\nabla_a \nabla_b Z^d - \nabla_b \nabla_a Z^d).$$

Wir bringen jeweils ein Vektorfeld unter die Ableitung und erhalten dann

$$X^a Y^b Z^c R_{abc}{}^d = X^a \nabla_a (Y^b \nabla_b Z^d) - Y^b \nabla_b (X^a \nabla_a Z^d) - (X^a \nabla_a Y^b - Y^a \nabla_a X^b) \nabla_b Z^d.$$

Die rechte Seite wird auch oft in der Form

$$[\nabla_X \nabla_Y - \nabla_Y \nabla_X - \nabla_{[X,Y]}] Z \quad (\text{A.4.2})$$

geschrieben, wobei $[X, Y]$ der *Kommutator* der beiden Vektorfelder X und Y ist (vgl. Übung)

Wir spezialisieren jetzt auf den Fall $X = \partial_i$ und $Y = \partial_k$. Dann bekommen wir zuerst

$$[\partial_i, \partial_k] = 0,$$

so dass der dritte Term in (A.4.2) verschwindet. Mit (A.4.1) erhalten wir

$$\begin{aligned}
[\nabla_{\partial_i} \nabla_{\partial_k} - \nabla_{\partial_k} \nabla_{\partial_i}] \partial_l &= \nabla_{\partial_i} \left(\sum_j \Gamma_{kl}^j \partial_j \right) - \nabla_{\partial_k} \left(\sum_j \Gamma_{il}^j \partial_j \right) \\
&= \sum_j \left[\frac{\partial \Gamma_{kl}^j}{\partial x^i} - \frac{\partial \Gamma_{il}^j}{\partial x^k} \right] \partial_j + \sum_{jm} \left[\Gamma_{kl}^j \Gamma_{ij}^m - \Gamma_{il}^j \Gamma_{kj}^m \right] \partial_m \\
&= \sum_m \left(\frac{\partial \Gamma_{kl}^m}{\partial x^i} - \frac{\partial \Gamma_{il}^m}{\partial x^k} + \sum_j \left[\Gamma_{kl}^j \Gamma_{ij}^m - \Gamma_{il}^j \Gamma_{kj}^m \right] \right) \partial_m
\end{aligned}$$

Damit ergeben sich schließlich die Komponenten des Riemann-Tensors aus den Zusammenhangskoeffizienten wie folgt

$$R_{ikl}{}^m = \frac{\partial \Gamma_{kl}^m}{\partial x^i} - \frac{\partial \Gamma_{il}^m}{\partial x^k} + \sum_j \left[\Gamma_{kl}^j \Gamma_{ij}^m - \Gamma_{il}^j \Gamma_{kj}^m \right]$$